

WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 36 | 30 August to 5 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

6

Industry / FDA items

3

Media highlights

4

Key opinion leaders

GOVERNANCE TAKEAWAY OF THE WEEK

This week the monitoring layer itself came under the microscope. The ACR's Assess-AI registry, the first national imaging AI registry, monitors AI outputs by having an LLM extract findings from signed reports, and its authors admit the extractor's own accuracy is not yet independently validated. Read that next to CIDER, where an extraction LLM was held to per-field accuracy and repeated-run reproducibility standards, and next to a TI-RADS study where a prompt change alone moved scoring accuracy by 28 points. The action for a radiology leader: inventory every place an LLM silently turns text into data in your department, including inside your monitoring tools, and give each one the same acceptance test you would give a diagnostic model, with the prompt version-locked and test-retest agreement measured before its output counts as evidence.

SECTION 0

EXECUTIVE SUMMARY

Key trends, week of 30 August to 5 September 2026

[MONITORING] THE ACR JUST SHOWED WHAT POST-MARKET MONITORING LOOKS LIKE AT NATIONAL SCALE

Schmidt K et al. (JACR) describe Assess-AI, the first national imaging AI registry: LLM prompts extract findings from signed reports and pair them with AI outputs across nine use cases, with development-cohort agreement of 0.985 for ICH and 0.997 for PE. The authors say plainly that these were tuning cohorts, not independent validation. The architecture is the takeaway: the signed report as recurring ground truth.

[WORKFLOW] A 90.6% TURNAROUND REDUCTION, MEASURED PROSPECTIVELY, WITH THE RIGHT FRAMING

Sridharan S et al. (JMIR) ran a prospective crossover on 1054 chest radiographs with 8 radiologists: AI-triaged worklists plus report generation cut median report time from 2 to 0.53 minutes and mean turnaround from 876 to 82 minutes across all urgency categories. The authors call AI workflow infrastructure, not a diagnostic replacement; missed-finding rates were not part of the endpoint set.

[VALIDATION] AN LLM EXTRACTOR IS ONLY A DATABASE IF IT GIVES THE SAME ANSWER TWICE

Posta M et al. (JMIR) validated CIDER, a locally deployed LLM pipeline, on 2073 Hungarian histopathology reports: exact-match accuracy ranged from 99.5% for sex down to 78.1% for tumor size, with reproducibility tested across 3 independent runs and temperatures 0 to 2.0. Per-field accuracy plus test-retest agreement is the validation template for every LLM that feeds a registry.

[DATA SHARING] DE-IDENTIFICATION BY INPAINTING PUTS SYNTHETIC PIXELS IN YOUR SHARED DATA

Kline A et al. (JHIM) redact burned-in PHI (F1 0.891, recall 0.912) and inpaint the gap with Stable Diffusion 2; downstream segmentation holds at Dice 0.936 to 0.959. Recall of 0.912 still leaves residual PHI to audit, and inpainted regions must be labelled so downstream users know which voxels are real.

[WATCHPOINT] THE MONITOR IS A MODEL TOO

Three of this week's papers stack into one lesson: the ACR's national registry monitors imaging AI using LLM extraction, CIDER shows how such an extractor should be validated (per-field accuracy, repeated runs, fixed temperature), and the TI-RADS study shows a prompt change alone can move accuracy by 28 points. If your post-market monitoring depends on an LLM reading reports, that LLM needs its own acceptance test, its own version control, and its own place in the risk file.

SECTION 1

PEER-REVIEWED PAPERS

Source: PubMed | Verified indexing | No preprints | 30 August to 5 September 2026

Streamlined TI-RADS Reporting From Bilingual Speech Recognition to LLM-Assisted Conclusion Generation Using a Two-Stage Workflow.

Han T, et al. | *Academic radiology* | 2026-09-01

Before you connect speech recognition and an LLM into structured reporting, treat the prompt as a configuration item that swings accuracy by nearly thirty points. In this two-stage feasibility study on thyroid ultrasound, an LLM-based speech recognition model beat the commercial one on error rate and correction time, and TI-RADS scoring accuracy depended heavily on prompting strategy, moving from roughly half correct with minimal prompting to about four in five with detailed guideline-based prompts. The authors conclude that deterministic post-processing remained essential and prospective validation is required before deployment.

Key metrics: 149 thyroid ultrasound cases. LLM-based ASR vs commercial ASR: word error rate 1.0 (IQR 0 to 2.0) vs 3.0 (IQR 2.0 to 3.0), $p < 0.001$; correction time 6.38 vs 11.39 seconds, $p < 0.001$; radiologist preference 70.5% and 73.7%. ACR TI-RADS point accuracy rose from 50.3% (minimal prompting) to 78.5% (detailed prompting), $p < 0.001$; K-TIRADS category accuracy from 58.4% to 79.9%, $p < 0.001$.

Clinical relevance: The 28-point accuracy swing between prompting strategies is the governance finding: the same model with a different prompt is a different device in everything but name. For your AI committee this changes acceptance testing: version-lock the prompt alongside the model, re-run the acceptance test whenever either changes, and keep deterministic rule-based checks on top of LLM output, consistent with Article 15 accuracy and robustness.

PMID 42680653 DOI 10.1016/j.acra.2026.08.009

Multi-task Deep Learning via UNETR for PSMA PET/CT Image for Prostate Cancer: Simultaneous Assessment of Prostate Cancer Disease Burden, Treatment Selection and Survival Outcomes.

Tun HM, et al. | *Journal of imaging informatics in medicine* | 2026-08-31

Kept on scientific merit, including the negative part: this multi-task framework for PSMA PET/CT tried to do lesion segmentation, treatment stratification and survival prediction from a single model, and the comparison came out against the elegant solution. A simple CNN reached an AUC of 0.91 for treatment classification, while the multi-task UNETR managed 0.561, buying interpretability through segmentation-based tumor quantification at a steep cost in discrimination. The authors report this plainly and follow TRIPOD-AI, which is more than many stronger-looking papers do.

Key metrics: 212 prostate cancer cases, 478 observations (2018 to 2024), single center; 205 matched PSMA PET/CT scans from 115 patients analyzed. Simple CNN treatment classification AUC 0.91, with active surveillance AUC 0.99 and chemotherapy AUC 0.97. Multi-task UNETR weighted AUC 0.561. Deceased patients showed higher Dice (mean 0.556, SD 0.065) than survivors (mean 0.324, SD 0.165).

Clinical relevance: PSMA PET/CT drives real treatment decisions across the prostate cancer pathway, and a model that stratifies CHARTED and LATITUDE eligibility from imaging alone would be clinically useful. This one is not there: single-center, modest cohort, and a multi-task architecture that underperforms a simple baseline on the task that matters. Its value is as a clearly reported design lesson that task-sharing does not come free.

PMID 42675275 DOI 10.1007/s10278-026-02224-3

Impact of AI-Triaged Worklists and AI-Assisted Report Generation on Radiology Turnaround Times: Prospective Real-World Study.

Sridharan S, et al. | *Journal of medical Internet research* | 2026-08-31

Before you pilot worklist triage or report generation, fix your baseline: this is one of the few prospective real-world measurements of what AI does to radiology turnaround, and the effect sizes are large enough to reorganize a service around. Eight board-certified radiologists read chest radiographs in two sessions, standard first-in first-out versus an AI-triaged worklist with integrated report generation, separated by a 4-week washout. The gains concentrate exactly where you would want them: normal studies cleared fastest, and turnaround fell across every urgency category including critical studies.

Key metrics: 1054 chest radiographs, single-center prospective paired study, single-sequence crossover, 8 radiologists. Median report generation time 2 minutes (IQR 1 to 4) unaided vs 0.53 minutes (IQR 0.22 to 1.12) AI-assisted, a 73.3% reduction; normal radiographs fell from 2 to 0.2 minutes, a 90% reduction. Mean turnaround time 876.21 minutes (SD 1014.20) vs 82.25 minutes (SD 83.38), a 90.6% reduction; all urgency categories $p < .001$.

Clinical relevance: The authors frame AI as workflow infrastructure rather than a diagnostic replacement, and that is the right lens for reading a 90.6% turnaround reduction: part of the effect is the triage logic itself replacing first-in first-out, not the model reading images better than humans. Single-center, single-sequence crossover, chest radiographs only. For your AI committee this changes acceptance testing: an operational AI pilot needs pre-declared efficiency endpoints measured per urgency category against your own baseline, plus ongoing monitoring that the time saved on normals is not being paid for in missed findings, which this study did not measure.

PMID 42673581 DOI 10.2196/92181

AI Monitoring AI with LLMs: The American College of Radiology's Imaging AI Registry.

Schmidt K, et al. | *Journal of the American College of Radiology* | 2026-09-03

Before you build local AI monitoring from scratch, look at how the ACR is doing it nationally. Assess-AI, the first national imaging AI registry, uses large language model prompts to extract clinically relevant findings from radiology reports and pair them with AI outputs, giving report-anchored performance monitoring at scale. The authors are careful about what this is not: the tuning cohorts that produced the headline agreement numbers also informed prompt refinement and label construction, so they are not independent validation sets, and clinical utility remains unestablished.

Key metrics: Prompts developed for nine use cases via AWS Bedrock LLMs, with consistency checked across 10 repeated runs per cohort. Final-prompt agreement with hybrid report-derived reference labels in Stage 2 development cohorts: 0.985 for intracranial hemorrhage, 0.997 for pulmonary embolism. The paper states these cohorts were not independent validation sets.

Clinical relevance: This is the reference architecture for post-market monitoring in radiology: the signed report becomes the recurring ground truth, and an LLM layer scales the comparison. The honest caveat matters as much as the design, because the extraction layer is itself a model whose accuracy on your local reports is unproven. For your AI committee this changes monitoring: registry participation or a local equivalent gives you the Article 72 post-market monitoring evidence stream, but the LLM extractor must pass its own local acceptance test before its concordance numbers count as evidence.

PMID 42692226 DOI 10.1016/j.jacr.2026.08.022

A Secure, Scalable Large Language Model-Based System (CIDER) for High-Throughput Clinical Data Extraction From Medical Reports: Retrospective Validation Study.

Posta M, et al. | *Journal of medical Internet research* | 2026-09-02

Before you let an LLM populate a registry or research database, demand the two numbers this team reported and most vendors do not: field-level accuracy against an expert-curated reference and reproducibility across independent runs. CIDER is an institutionally deployable extraction pipeline built on a locally hosted Qwen3-VL-32B model, validated on 2073 real-world Hungarian-language histopathology reports, a deliberately hard non-English setting. Extraction quality was high but visibly uneven across fields, which is exactly the profile you should expect and test for.

Key metrics: 2073 Hungarian histopathology reports, seven variables, expert-curated reference standard. Exact-match accuracy at temperature 0.1: sex 99.5%, surgery year 98.1%, organ 95.8%, T stage 95.6%, N stage 92.4%, histology 87.5%, tumor size 78.1%. Weighted F1 0.87 to 0.995; Cohen kappa 0.85 to 0.99. Tumor size within 5 mm in 83.5% to 85.3% and within 10 mm in 87.3% to 88.7% of cases. Robustness assessed across temperatures 0 to 2.0; reproducibility across 3 independent runs.

Clinical relevance: Locked-down local deployment answers the data security objection, and the Hungarian setting shows serious multilingual capability, but the operational lesson is the validation design: per-field accuracy plus repeated-run reproducibility at fixed temperature. A model that gives different answers on different runs is not a database. For your AI committee this changes validation: any LLM extraction tool, including the ones inside monitoring systems, should pass a per-field accuracy threshold and a test-retest agreement threshold on your own documents before its output is treated as data, in line with Article 15 accuracy and robustness.

PMID 42686197 DOI 10.2196/95780

Automatic Kidney Image Segmentation During Robot-Assisted Partial Nephrectomy Using a Deep Learning Model Based on a Multiannotator Dataset: Model Development and Validation Study.

Margue G, et al. | *JMIR medical informatics* | 2026-09-02

Kept on scientific merit: this is the unglamorous groundwork real-time surgical augmented reality actually depends on. From 131 robot-assisted partial nephrectomy procedures, the team built a 48,000-frame dataset, trained nine annotators of mixed backgrounds under a structured protocol, and used 12,546 annotated images to train a kidney segmentation network. The honest headline is the gap: annotators agreed with the expert reference at Dice 0.91 to 0.95, while the model reached 0.75, dropping further during tumor resection and reconstruction when the scene gets visually complex.

Key metrics: 131 RAPN procedures (2022 to 2024), 160 tumors, 48,000 extracted frames. Interannotator agreement on a 454-image subset: median Dice 0.91 to 0.95, sensitivity above 0.89. AlbuNet-34 trained on 12,546 images, validated on 3137: mean Dice 0.75 (SD 0.23), sensitivity 0.71 (SD 0.24), with significant performance variation across surgical phases.

Clinical relevance: Intraoperative segmentation is a harder problem than radiology segmentation because the anatomy moves, bleeds and deforms, and this study quantifies exactly how far current models are from the expert level a surgeon would need to trust an overlay. The phase-stratified reporting is the useful part: performance collapse during resection, the moment AR guidance would matter most, is the number that defines the remaining work.

PMID 42685340 DOI 10.2196/82540

SlideChat is a multimodal generative artificial intelligence assistant for whole-slide computational pathology across cancer types.

Chen Y, et al. | *Nature cancer* | 2026-09-01

Kept on scientific merit: SlideChat moves multimodal pathology AI from patch-level snippets to whole-slide reasoning. It couples patch-level and slide-level pathology encoders to a pretrained large language model and was trained on 274,233 multimodal instruction samples to link gigapixel whole-slide images to diagnostic reports and answer clinical queries across cancer types. The evaluation is unusually broad for this class of system, spanning closed-ended and open-ended questions plus report generation over five cohorts.

Key metrics: Trained on 274,233 multimodal instruction samples. Evaluated on 8,836 closed-ended questions, 129 open-ended questions and 3,149 whole-slide image reports from five cohorts spanning 31 cancer types. Outperformed leading baselines by 19.1% in closed-ended accuracy and 7.7% in report-generation METEOR score; highest expert ratings across five dimensions in open-ended answering. The abstract reports comparative gains rather than absolute accuracy values.

Clinical relevance: Whole-slide context is what pathologists actually use, and an assistant that reasons at slide level rather than patch level addresses the main structural criticism of computational pathology to date. The evaluation is benchmark and expert-rating based; the abstract makes no claim about prospective diagnostic performance or patient outcomes, and the gains are reported relative to baselines rather than as absolute figures.

PMID 42680939 DOI 10.1038/s43018-026-01220-4

From Redaction to Restoration: Deep Learning for Medical Image Deidentification and Reconstruction.

Kline A, et al. | *Journal of imaging informatics in medicine* | 2026-08-31

Before you release or pool imaging data, decide what happens to the pixels you redact. This framework detects burned-in protected health information, redacts it, and then uses a latent-diffusion model to inpaint the redacted regions with anatomically plausible content, so that downstream analysis tools receive intact-looking volumes instead of black boxes. The privacy performance and the preserved downstream utility are both quantified, which makes this a rare de-identification paper you can actually evaluate.

Key metrics: PHI region detection F1 0.891 (SD 0.037), recall 0.912 (SD 0.053), precision 0.875 (SD 0.058). Downstream segmentation remained stable across inpainting strategies: Dice 0.936 to 0.948 on one dataset and 0.955 to 0.959 on another. Pipeline combines CRNN-based redaction with Stable Diffusion 2 inpainting.

Clinical relevance: A recall of 0.912 means residual PHI in roughly one in eleven PHI-containing regions, so this is a component of a de-identification process with human QA, not a replacement for one. And inpainting means shared volumes contain synthetic pixels: harmless for the tested segmentation tasks, but a fact that must travel with the data. For your AI committee this changes data governance under Article 10: any release pipeline using generative restoration needs a residual-PHI audit step and explicit labelling of inpainted regions in the dataset documentation, so downstream users know which voxels are real.

PMID 42675278 DOI 10.1007/s10278-026-02241-2

SECTION 2

INDUSTRY & REGULATION

Sources: Official registers | Press releases | FDA databases

[FDA] Heartvue.ai receives FDA 510(k) clearance for Heartvue.Proton cardiac MRI analysis software, with an authorized Predetermined Change Control Plan

Heartvue.ai | 2 September 2026

Cleared for: viewing, post-processing and quantitative evaluation of cardiovascular MR images, automating 2D linear dimensions, ventricular volumes and ejection fraction, and blood flow quantification, with every result under physician review (K260811, decision date 26 August). Evidence level: vendor and consultancy announcement; the release cites the clinical validation study designed for the submission but no published peer-reviewed performance data. The notable feature is the FDA-authorized PCCP, allowing model updates within pre-specified bounds without a new 510(k). Questions for the vendor: the quantitative agreement with expert manual measurement by sequence and scanner; which model changes the PCCP covers and how you will be notified when one ships; whether per-version performance reports are provided. A PCCP shifts change control onto the deployer: log the software version with every study so you can segment your own QA by model version.

Source: [Business Wire via The AI Journal](#)**[REGULATION] Vara secures the first CE mark for an autonomous AI breast cancer screening triage system, Class IIb under MDR**

Vara | 2 September 2026

Certified for: autonomous triage in organised population breast screening, stratifying mammograms into normal versus needing radiologist assessment, as a Class IIb device under MDR; the first authorization of its kind globally. Evidence level: strong by imaging AI standards, anchored on the prospective PRAIM study (NCT04778670) of 461,818 asymptomatic women aged 50 to 69 across 12 German sites, with detection reported noninferior and statistically superior versus control, plus the ATMON real-time supervision system built on seven years of real-world monitoring. Questions for the vendor: the exact share of studies the system may close autonomously and on what threshold; sensitivity in the autonomously triaged normal stream, not just overall; how ATMON alarms reach your program and who owns the response; performance in your national screening population versus the German cohort. Autonomous closure of normals is the highest-stakes error direction in screening: define your interval-cancer audit on that stream before go-live.

Source: [Medical Device Network](#)**[FDA] Aidoc receives FDA Breakthrough Device Designation for First Read, AI drafting preliminary chest radiograph report text**

Aidoc | 4 September 2026

Designated for: AI analysis of chest radiographs generating preliminary radiology report text, with the designation (Q260882) granted for four life-threatening findings. This is not a clearance: First Read remains investigational and is not FDA cleared or approved, a distinction the coverage headlines blur. Evidence level: none published; Breakthrough Device Designation is a review-speed mechanism recognising unmet need, not a performance judgement. Questions for the vendor when it does reach market: which findings the cleared claim actually covers versus the broader drafting function; hallucination and omission rates per report against final signed reports; how automation bias in radiologists editing AI drafts was measured. Generative report drafting will need concordance monitoring of draft versus signed text from day one; a designation announcement is the moment to start designing that, not to start buying.

Source: [Medical Product Outsourcing](#)

[MARKET] Elucid raises an oversubscribed 55 million dollar Series D for AI coronary CT analysis, FFR-CT under FDA review

Elucid | 2 September 2026

The Boston-based coronary plaque analysis vendor announced 55 million dollars in new funding on 2 September, taking its total past 185 million, with participation from an unnamed publicly traded large medtech firm; it now counts four large public strategic investors. Funds are earmarked for platform expansion and its BioIntegrated FFR-CT product, currently under FDA review. Evidence level: financing news from the company; no new clinical data. For buyers in the coronary CT space the signal is competitive consolidation against Cleerly and Heartflow: ask any vendor in this category for head-to-head evidence against invasive FFR in a population like yours, per-lesion and per-patient, and how reimbursement actually lands in your setting before committing to a platform.

Source: [Radiology Business](#)

[MARKET] Scan.com closes 220 million dollars to run US imaging access as an API platform

Scan.com | 31 August 2026

Announced 31 August: 90 million in Series C equity plus 130 million in debt facilities for M&A and working capital, to expand a national platform combining imaging search, scheduling and results delivery through a single API, with AI matching referrals to availability, price and subspecialty. Evidence level: company announcement; the claimed 48-hour typical results turnaround is a vendor figure. This is network-layer infrastructure, not clinical AI, and it follows a different governance path: vendor management, API security and service-level enforcement rather than clinical validation. Questions before integrating: how radiologist subspecialty routing is verified; what quality oversight applies to a provider network that will change rapidly under an M&A strategy; who is accountable when an automated scheduling or routing decision delays a time-sensitive study.

Source: [MarketScale](#)

[PARTNERSHIP] CARPL.ai and GenServe partner to connect radiology AI outputs to autonomous patient outreach and follow-up workflows

CARPL.ai / GenServe.AI | 4 September 2026

Announced 4 September: GenServe's voice AI and autonomous care workflows (patient outreach, scheduling, referrals, follow-up completion) become available through CARPL's vendor-neutral radiology AI platform, aiming to act on AI-generated findings such as follow-up imaging recommendations. Evidence level: partnership announcement only; no outcome data on follow-up completion rates. The governance point is scope creep: an imaging AI finding that triggers an autonomous phone call to a patient has left the radiology QA perimeter and entered care delivery. Questions to ask: who clinically reviews an AI finding before outreach is triggered; how errors propagate and are recalled once a patient has been contacted; what audit trail links each outreach action back to the model version and finding that caused it.

Source: [CB Herald](#)

SECTION 3

POST-MARKET SIGNAL

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

Vara's ATMON: a CE-certified real-time supervision layer for autonomous screening AI, licensable on its own

Medical Device Network | 2 September 2026

Inside Vara's autonomous breast screening CE mark announced this week sits the more transferable post-market story: certification rested on ATMON, a purpose-built system for real-time supervision of the autonomous AI, grounded in seven years of continuous real-world monitoring, and Vara says it will license ATMON independently so providers could wrap it around other autonomous screening tools. This is post-market monitoring promoted from an obligation to a certified product component: the regulator accepted autonomy only with a live safety net attached. For deployers the precedent is concrete: if a vendor proposes any autonomous operation, ask what the equivalent of ATMON is in their offer, what it monitors in real time, and who receives its alarms.

Link: <https://www.medicaldevice-network.com/news/vara-ce-mark-ai-autonomous-breast-can...>

One year before, one year after: Erasmus MC measures what nodule AI actually did to reporting times in production

AuntMinnie / Radiology | 1 September 2026

A study published 1 September in Radiology is a model of post-deployment measurement: Erasmus MC compared 19,190 chest CT exams in the year before deploying a commercial nodule tool (Veye Lung Nodules) with 20,133 in the year after, using PACS-derived reporting time as the endpoint. Median reporting time fell 14.6% (21.3 to 18.2 minutes), with the largest gains for ECG-gated thoracic exams (-41.1%) and thoracic subspecialists (-25.0%), while emergency department exams got 7.1% slower. That last number is the monitoring lesson: aggregate benefit coexisted with a measurable subgroup harm signal, visible only because the department segmented its own routine data. Exploratory modelling put the saving near 0.5 FTE at institutional volumes.

Link: <https://www.auntminnie.com/clinical-news/ct/article/15833859/ai-cuts-chest-ct-re...>

SECTION 4

PLAYBOOK SNIPPET

One concrete step your AI committee can take this month

EU AI Act Article 15 (accuracy, robustness and cybersecurity)

Step: Add reproducibility as a separate acceptance test for any LLM-based tool: run the same 30 of your own documents or cases through it twice, days apart, at the vendor's production settings. Report test-retest agreement alongside accuracy and set a threshold below which the tool does not go live, with the prompt and model version recorded in the acceptance file.

Why: A model that answers differently on Tuesday is not deployable however accurate it was on Monday, and accuracy testing alone will never surface this.

Worked example: Posta M et al. validated the CIDER extraction pipeline on 2073 Hungarian histopathology reports with exactly this design: robustness assessed across temperatures 0 to 2.0 and technical reproducibility across 3 independent runs at temperature 0.1, alongside per-field exact-match accuracy from 99.5% (sex) down to 78.1% (tumor size). J Med Internet Res, PMID 42686197, DOI 10.2196/95780.

SECTION 5

MEDIA HIGHLIGHTS

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

Neiman Institute: clinicians who self-interpret their own imaging order far more of it

AuntMinnie / ITN | 1 September 2026

Trade coverage of a JACR study of 24.2 million Medicare office visits: visits with clinicians who self-interpreted all their ordered imaging had 188% higher odds of imaging overall than visits with non-self-interpreters, significant across every modality (54% higher for MRI up to 251% for ultrasound), and practices referring all imaging to radiologists showed 5% to 41% lower odds across most modalities. The authors estimate office-based imaging could fall about 17% without self-interpretation. ACR leadership tied the finding to support for appropriate-use-criteria legislation. For AI governance the relevance is upstream of any algorithm: utilization incentives shape the case mix every deployed model sees.

Link: <https://www.auntminnie.com/practice-management/news/15833771/selfinterpreting-by...>

Medicare's 137.53 dollar NTAP for CT triage AI: what the add-on payment actually requires

MarketScale | 30 August 2026

An analysis piece unpacking the CMS new technology add-on payment of up to 137.53 dollars per eligible inpatient case for Aidoc's CARE Multi-Triage CT Body, billable from 1 October for fiscal 2027 under a three-year window. The useful part for deployers is the operational fine print: payment depends on documented use identified through procedure coding, meaning PACS workflow, coding and analytics teams must reconcile AI-ran versus AI-billed each month, and the piece argues procurement should now write post-go-live monitoring and evidence-collection requirements directly into AI contracts. Reimbursement arriving with documentation obligations attached is the pattern to expect.

Link: <https://www.marketscale.com/industries/healthcare/medicare-sets-13753-maximum-ad...>

Rad AI interviews 12 radiologists on what is actually breaking in health systems

Rad AI | 1 September 2026

A vendor-published but unusually grounded series drawn from interviews with 12 radiologists across career stages and practice settings, published 1 September. The systemic numbers frame the AI conversation: 67% of practices understaffed amid rising volumes, real per-patient reimbursement down 25% between 2005 and 2021, and attrition surfacing downstream as ED length of stay and stroke door-to-needle slippage. The recurring practitioner view is that the hard problem is operationalizing AI within revenue cycle and peer-learning processes, not the algorithms, and that the radiologist's consultative role is what deployment plans keep undercounting. Source is a vendor blog; read it as field testimony, not evidence.

Link: <https://www.radai.com/blogs/the-radiology-conversation-health-systems-arent-havi...>

SECTION 6

PROFESSIONAL SOCIETIES

SIIM | ACR | RCR | Official publications and event coverage

ACR files comments on the CY2027 HOPPS proposed rule, including Software as a Medical Service payment

2 September 2026

ACR submitted comments to CMS this week on the calendar year 2027 Hospital Outpatient Prospective Payment System proposed rule, cautioning against the site-neutral payment proposal for imaging without contrast in off-campus departments, urging exclusion of preventive imaging and a payment floor, and providing recommendations on Software as a Medical Service, APC assignments and diagnostic radiopharmaceutical payment. The SaMS thread is the one imaging AI vendors and deployers should track: how CMS categorizes and pays for algorithmic services in the outpatient setting is being negotiated now, ahead of the final rule this fall.

Source: <https://www.acr.org/News-and-Publications/2026/site-neutral-payment-in-hopps-comments>

ACR announces three new draft QCDR measures for 2027 MIPS

3 September 2026

The ACR Qualified Clinical Data Registry team announced on 3 September that it intends to submit three new QCDR measures for CMS consideration for the 2027 MIPS performance year, alongside twelve existing QCDR measures and the new Diagnostic Radiology and Interventional Radiology MIPS Value Pathways, and continues streamlining reporting for eCQM 494 on excessive radiation dose or inadequate CT image quality, currently the only outcome measure in the Diagnostic Radiology MVP. Quality measurement infrastructure of this kind is where AI-derived metrics will eventually have to land to affect payment.

Source: <https://www.acr.org/News-and-Publications/2026/draft-measures-for-2027-mips>

RSNA's inaugural Across Africa: Advancing Cancer Imaging course runs in Cape Town

4 September 2026

The first event of RSNA's Across Africa program is taking place 4 to 5 September at the Biomedical Research Institute in Cape Town, under the society's Global Health Initiative with sponsorship from GE HealthCare. The two-day programme combines cancer imaging clinical sessions with capacity-building case studies from Egypt, Tanzania, Uganda and Kenya, and closes with a session on AI innovation in global cancer care, including digital health and AI for oncology in sub-Saharan Africa. A second half-day session follows at RSNA 2026 in Chicago on 2 December.

Source: <https://www.rsna.org/news/2026/july/cancer-imaging-education-cape-town>

SECTION 7

KEY OPINION LEADERS

LinkedIn | Public statements | Week of 30 August to 5 September 2026

Curt Langlotz

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

Posted 4 September, sharing Radiology Business coverage of the first regulatory approval for a breast screening triage tool that closes normal studies without radiologist review, with a two-word verdict: 'It's happening...'. Coming from one of the field's most careful voices, the brevity is the message: autonomous operation in screening has moved from panel debate to regulatory fact, and the community conversation now shifts from whether to under what supervision. The post drew over a hundred reactions and a repost wave within a day.

Source: [LinkedIn post](#), 4 September 2026

Woojin Kim

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

Posted 3 September previewing his three roles at KCR 2026 in Goyang (9 to 12 September), including a lecture 'From Foundation Models to Agentic AI: Building the Future of MSK Radiology', and posted 4 September announcing his eight AI-generated entries in RSNA's Art of Imaging AI contest, riffing on classic radiologic signs from the Scottie dog to crazy paving. He also amplified ACR Board Chair Christoph Wald's post on the newly published 'Detection, Diagnostic, and Triage AI Buyer's Guide: Questions that Matter', a risk-based evaluation framework for calculating the AI errors that matter to a practice's own population.

Source: [LinkedIn posts](#), 3 and 4 September 2026

Pranav Rajpurkar

Associate Professor, Harvard; co-founder, a2z Radiology AI

Posted 5 September on the ReXGroundingCT Challenge at MICCAI 2026: 36 teams competing, top score climbed from 0.332 to 0.367 Dice, ten days remaining, with the leaderboard led by Gachon University, DEEPNOID and the Netherlands Cancer Institute. The benchmark matters for the report-generation debate: grounding, making a model show which pixels support each sentence of a generated finding, is the technical capability that report-drafting AI will need before its output can be audited rather than merely read.

Source: [LinkedIn post](#), 5 September 2026

Amine Korchi

Radiologist, healthtech investor; Sifted Top 25 Expert, Imaging Wire Top 10 Radiology AI KOL

Posted 4 September announcing a move into long-form video: after years of using LinkedIn as his main channel for radiology, healthtech and investing commentary, he is launching a YouTube presence for ideas that 'need more room' for context and nuance beyond the headline. A small signal of a broader shift: the imaging AI conversation is migrating toward formats where evidence and reasoning can be unpacked, not just announced.

Source: [LinkedIn post](#), 4 September 2026

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 30 August to 5 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.