

WEEKLY AI INTELLIGENCE REPORT

Radiology and Medical Imaging

Week 37 | 6 to 12 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

3

Key opinion leaders

GOVERNANCE TAKEAWAY OF THE WEEK

This week hands a radiology leader a clean two-way split in evidence maturity. AI assistance for image interpretation now has trial-grade support: a five-centre randomised crossover in Lancet Digital Health shows an ultrasound aid raising sonographer sensitivity for fetal intracranial malformations by 0.087 while formally protecting specificity. LLM report generation has nothing comparable: a systematic review of 101 studies found not one at low risk of bias and safety data too heterogeneous to pool. The action: split your AI pipeline into these two evidence classes. For interpretation aids, require reader-performance evidence of the kind that now demonstrably exists. For LLM drafting, treat every pilot as building the safety case from zero, on your own casemix, with your own error counts.

SECTION 0

EXECUTIVE SUMMARY

Key trends, week of 6 to 12 September 2026

[EVIDENCE] AN RCT JUST SET THE BAR FOR ULTRASOUND AI ASSISTANCE CLAIMS

Lei T et al. (Lancet Digital Health) randomised the sequence, not the patients: five centres, 1584 scans, sonographers reading with and without PAICS assistance against a masked expert reference. Sensitivity for targeted fetal intracranial malformations rose by 0.087 with specificity formally non-inferior. This is the design to demand from any vendor claiming their tool makes your readers better.

[REALITY CHECK] 101 STUDIES ON LLM REPORT GENERATION, NONE AT LOW RISK OF BIAS

Huang JL et al. (JMIR) reviewed the entire LLM report-generation literature to May 2026: 101 studies, zero judged at low risk of bias, safety and workflow outcomes too heterogeneous to pool. Acceptance rates approach radiologist reports while clinically significant errors persist. The safety case for LLM drafting is not in the literature; it has to be built locally.

[VALIDATION] YOUR FOUNDATION MODEL CHOICE IS A RESOLUTION CHOICE

Tayebi Arasteh S et al. (Communications Medicine) benchmarked DINOv3 on 816,183 chest radiographs: it beats DINOv2 only at 512 x 512 pixels, gains vanish in pediatric cohorts, 1024 x 1024 rarely helps, and a frozen 7B model lost to fully adapted mid-sized ones. Resolution, backbone and adaptation regime are validation parameters, not implementation details.

[INFORMATICS] THE UNCERTAIN CLASS IS WHAT MAKES REPORT MINING HONEST

Jo W et al. (JMIR Medical Informatics) classified 288,076 postoperative CT reports for recurrence and metastasis with an explicit uncertain category for hedged language, evaluated under simulated real-world class distributions, with clinicians reviewing under 1% of reports. Registry-scale surveillance from report text is workable when uncertainty is a label, not an error.

[WATCHPOINT] TWO EVIDENCE CLASSES, ONE PROCUREMENT DESK

This week shows the spread in what counts as evidence: a randomised crossover trial for ultrasound AI assistance at one end, and a 101-study review of LLM report generation with no low risk-of-bias study at the other. Detection and assistance aids can now be held to reader-performance evidence; LLM drafting cannot yet. Write your acceptance tests accordingly: demand the trial-grade design where it exists, and budget for building the safety case yourself where it does not.

SECTION 1

PEER-REVIEWED PAPERS

Source: PubMed | Verified indexing | No preprints | 6 to 12 September 2026

Evaluating AI-assisted detection of fetal intracranial malformations in prenatal ultrasound practice: a multicentre, self-crossover, randomised controlled trial in China.

Lei T, et al. | *The Lancet. Digital health* | 2026-09-10

Before you accept standalone AUC as evidence for an ultrasound AI aid, ask the vendor for this design: a multicentre, self-crossover randomised controlled trial at five Chinese centres measured what the PAICS system does to sonographer performance, not what the model does alone. Sonographers with 3 to less than 8 years of experience scanned high-risk pregnancies with and without AI assistance, with a 4-week washout, concealed allocation and an expert video-review panel as reference standard.

Key metrics: 1584 scans completed across five centres. Fetal-based sensitivity for ten targeted intracranial malformations increased by 0.087 (95% CI 0.029 to 0.147, p superiority < 0.0001); specificity non-inferior (difference 0.009, 95% CI -0.006 to 0.023, p non-inferiority < 0.0001). Malformation-targeted sensitivity increased by 0.118 (95% CI 0.054 to 0.181). No adverse events reported. Registration ChiCTR2200063424.

Clinical relevance: This is the strongest evidence class an imaging AI assistance claim can carry: a randomised crossover measuring the human-plus-AI effect against a masked reference standard, with specificity protected by a formal non-inferiority margin. The scope discipline matters too, as the trial separates malformations inside the tool's recognition scope from anomalies beyond it. For your AI committee this changes acceptance testing: for assistance tools, require evidence on reader performance with the tool, not standalone accuracy, and record the tool's recognition scope in the deployment file; this is the operational content of EU AI Act Article 14 human oversight.

PMID 42722598 DOI 10.1016/j.landig.2026.101040

Abdominal landmark detection and classification of fetal growth conditions on obstetric ultrasound: a gestational age-conditioned multi-task network.

Tong T, et al. | *Abdominal radiology (New York)* | 2026-09-10

A multi-task obstetric ultrasound network shows what disciplined clinical model development looks like. GA-MTNet conditions image features on gestational age, then jointly detects four abdominal landmarks, classifies five sonographic patterns linked to fetal growth conditions and regresses biometry, trained on 6,750 images from two centres and tested on a third centre never seen in development.

Key metrics: 9,046 B-mode images from 2,712 pregnant women (18 to 34 weeks) across three institutions. Internal macro-AUC 0.863, accuracy 84.7%, landmark mAP@0.5 0.791. External test (1,296 images, 434 patients): AUC 0.841, accuracy 82.3%, a 2.2-point AUC drop. Gestational age conditioning added 3.1 percentage points ($p = 0.003$). Calibration ECE 0.047.

Clinical relevance: The clinically interesting part is the design: gestational age as an explicit model input rather than a confounder, a held-out external site, ablations that quantify each component, and reported calibration. The continual learning mechanism (Monte Carlo Dropout uncertainty, elastic weight consolidation, low-rank adaptation) is proposed for post-deployment refinement but is evaluated here in development conditions, not in service.

PMID 42720764 DOI 10.1007/s00261-026-05771-z

Resolution-dependent self-supervised transfer in chest radiograph classification.

Tayebi Arasteh S, et al. | [Communications medicine](#) | 2026-09-09

Before you standardise on a foundation model for chest radiography, fix the input resolution first. This benchmark of DINOv3 against DINOv2 and supervised ImageNet initialisation across seven datasets finds the newest self-supervised model wins only at 512 x 512 pixels, mainly with a ConvNeXt-B backbone, and mainly for small focal and boundary-dependent abnormalities.

Key metrics: 816,183 radiographs across seven pediatric and adult datasets; primary outcome mean AUROC. DINOv3 not consistently better than DINOv2 at 224 x 224 but strongest initialisation at 512 x 512. Scaling to 1024 x 1024 rarely improved performance at markedly higher compute. Pediatric cohort showed no significant benefit from DINOv3, resolution or backbone. Frozen DINOv3-7B features underperformed fully adapted 86 to 89M-parameter models.

Clinical relevance: The result cuts against two procurement reflexes: newer foundation model does not mean better transfer, and bigger frozen backbones lost to fully adapted mid-sized ones. Benefits also did not carry from adult to pediatric cohorts, so population transfer must be tested, not assumed. For your AI committee this changes validation: acceptance tests for anything built on a foundation model should fix and record resolution, backbone and adaptation regime as tested configurations under EU AI Act Article 15 accuracy and robustness, and re-run when any of them changes.

PMID 42717246 DOI 10.1038/s43856-026-01897-9

Automated Extraction of Postoperative Cancer Recurrence and Metastasis From Computed Tomography (CT) Reports: Semisupervised Deep Learning Study.

Jo W, et al. | [JMIR medical informatics](#) | 2026-09-08

Cancer surveillance at registry scale needs the uncertainty category, not just positive and negative. This semisupervised framework classifies postoperative CT reports for recurrence and metastasis into positive, negative and uncertain, explicitly modelling the hedged language (cannot exclude recurrence, possibly metastatic) that defeats binary extraction, with clinicians reviewing under 1% of reports across three human-in-the-loop cycles.

Key metrics: 288,076 postoperative CT reports from 86,083 cancer surgery patients (2014 to 2021); 17,846 unique reports for recurrence and 63,766 for metastasis after deduplication. Under simulated real-world class distributions, PubMedBERT reached 92.58% accuracy for multiclass recurrence and MedEmbed 93.25% for binary metastasis; the framework reached 97.33% (multiclass) and 99.33% (binary) for recurrence, 95.00% and 96.67% for metastasis.

Clinical relevance: For imaging informatics teams building outcome feeds from report text, the design points are the three-way label with an explicit uncertain class, evaluation under simulated real-world class distributions rather than balanced test sets, and a quantified human-review budget. Single-centre Korean data; transfer to other languages and reporting styles is untested.

PMID 42709940 DOI 10.2196/92937

Development and Validation of a Convolutional Neural Network Framework Based on Ultrasound Imaging for Multi-Classification of Superficial Soft Tissue Masses.

Wu M, et al. | [Academic radiology](#) | 2026-09-07

Superficial soft tissue masses are a recognised weak spot of subjective ultrasound reading, and this four-institution study builds a staged classifier for them. ST-USNet chains four sub-models: benign versus malignant, malignant subtyping (sarcoma, lymphoma, metastatic carcinoma), lymphoma aggressiveness, and an exploratory head for metastatic origin.

Key metrics: Ultrasound images of 3,168 patients (median age 58, IQR 46 to 68) from four institutions, 2015 to 2024. Validation AUC 0.984 for benign versus malignant; 0.932, 0.909 and 0.922 for sarcoma, lymphoma and metastatic carcinoma; 0.951 for aggressive versus indolent lymphoma. Performance slightly lower on an independent test cohort. Preliminary reader study: 4 radiologists, 85 cases; the tool outperformed senior radiologists on one task ($p = 0.012$) and was comparable on the others.

Clinical relevance: The staged architecture mirrors the actual clinical decision sequence, and the authors are explicit that the metastatic-origin head is exploratory and that multi-center prospective validation and continuous updating are required before deployment. The drop from validation to independent test cohort is the number to watch when this class of tool reaches your desk.

PMID 42705922 DOI 10.1016/j.acra.2026.08.092

Effectiveness, Safety, and Workflow Burden of Large Language Model-Based Medical Report Generation: Systematic Review.

Huang JL, et al. | *Journal of medical Internet research* | 2026-09-09

Before you pilot LLM report drafting, read this systematic review of 101 studies for what is missing: not one study was judged at low risk of bias, and safety and workflow outcomes were so heterogeneous that no meta-analysis was possible. Effectiveness signals exist, including expert acceptance close to radiologist reports, but clinically significant errors persist in the same studies.

Key metrics: 101 studies, searched to 15 May 2026; chest x-ray the largest group (36 studies). Risk of bias: 15 moderate, 72 high, 14 serious, none low. In one chest x-ray study AI report acceptance was 70.5% versus 73.3% for radiologist reports, with false negatives 18.5% versus 17.8%. In a collaboration study AI reports were equivalent or preferred in 77.7% and 56.1% of cases across two datasets. A brain MRI study reduced reading time from 61 to 53 seconds.

Clinical relevance: The gap between benchmark performance and deployable evidence is the finding. Acceptance and preference are measured; omission and commission error rates, edit burden and oversight workload are sparsely and inconsistently reported, which is exactly the evidence a deployer needs. For your AI committee this changes both validation and monitoring: an LLM drafting pilot needs a local error taxonomy (omission, commission, clinically significant) counted on your own cases before go-live under EU AI Act Article 9 risk management, and a standing sample-based error audit after it, because the published literature will not carry the safety case for you.

PMID 42715525 DOI 10.2196/97007

Global Research Trends and Insights on AI Usage in Tuberculosis: Bibliometric Analysis.

He J, et al. | *JMIR medical informatics* | 2026-09-10

Where AI research for tuberculosis actually happens matters for anyone deploying imaging AI in high-burden settings. This bibliometric analysis of 1,300 articles from 2000 to 2025 maps a research-demand inversion: output concentrates in the United States while high-burden countries remain underrepresented relative to disease burden, and technical validation work markedly outweighs reported clinical application.

Key metrics: 1,300 articles from Web of Science Core Collection, 2000 to 2025, screened by two independent reviewers. Post-2018 annual growth: 29.7% in publications, 54.2% in citations. United States: 346 publications, 12,143 citations. Of the 15 most cited papers, 4 focused on imaging analysis, 2 on drug resistance prediction, 3 on drug discovery.

Clinical relevance: For TB programmes and their imaging partners the authors' conclusion is concrete: prioritise locally validated, implementation-ready tools over algorithmic refinements, and expect persistent gaps in cross-regional data sharing, atypical lesion recognition and socioeconomic factor integration. A map of where evidence is produced is also a map of where external validity is untested.

PMID 42721319 DOI 10.2196/93885

ParasiteNet: Deep Learning-Based Identification of Parasitic Helminth Eggs in Microscopic Images.

Batjargal D, et al. | *Journal of imaging informatics in medicine* | 2026-09-08

Helminth microscopy is a bottleneck exactly where expertise is scarcest, and this study benchmarks two architectures for automated egg classification. ParasiteNet compares a CNN plus ResNet50 hybrid against EfficientNetV2B0 across 12 helminth egg categories on annotated microscopic images.

Key metrics: 12,000 annotated images across 12 categories. EfficientNetV2B0: 96.92% classification accuracy, loss 0.0067. CNN plus ResNet50 hybrid: 95.25% accuracy, loss 0.1840.

Clinical relevance: The authors state the limits plainly: performance was demonstrated under controlled experimental conditions, and independent external datasets, diverse microscopy settings and prospective clinical studies are needed before real-world deployment. As a screening aid for low-resource settings the direction is promising; the field-condition evidence is the missing piece.

PMID 42711646 DOI 10.1007/s10278-026-02115-7

SECTION 2

INDUSTRY & REGULATION

Sources: Official registers | Press releases | FDA databases

[MARKET] Epsilon Health emerges from stealth with 27.6 million dollars for an AI-native radiology practice*Epsilon Health | 10 September 2026*

What it is: not a cleared device but a business model, a radiology practice that embeds its own AI into board-certified radiologists' workflow, announced 10 September with funding led by AlleyCorp. Claims: over 250,000 patients served, more than 2,500 studies daily, on track to read 1% of daily US x-rays in 2026. Evidence level: company statements, no published peer-reviewed performance data. Questions before referring or contracting: which functions are FDA-cleared devices versus internal practice-of-medicine tools; who holds clinical responsibility for AI-influenced reads; what per-study quality metrics are shared with clients; how the continuous model iteration on proprietary data is versioned and validated. The practice model shifts AI oversight from the buyer to the seller; your monitoring is the service-level agreement.

Source: [AuntMinnie](#) / [Business Wire](#)**[FDA] LUMA Vision announces expanded FDA clearance for VERAFFEYE cardiac imaging and navigation platform***LUMA Vision | 9 September 2026*

Cleared for: expanded capabilities of VERAFFEYE, a 4D intracardiac imaging, visualization and navigation platform for minimally invasive cardiac procedures, ahead of a commercial launch planned for early 2027. Evidence level: vendor announcement; the release names no reader study or published clinical performance data for the newly cleared capabilities. Questions for the vendor: the exact 510(k) intended-use wording and which capabilities are new; whether any AI or automated function influences guidance decisions and on what test data it was validated; what procedure-level outcome evidence exists versus imaging-quality claims. A platform clearance is not a claim of clinical benefit in any specific procedure; local evaluation still applies before first clinical use.

Source: [PR Newswire](#) / [BioSpace](#)**[MARKET] AEYE Health publishes its three pivotal trials for autonomous diabetic retinopathy screening***AEYE Health | 9 September 2026*

What happened: AEYE Health announced peer-reviewed publication in *Frontiers in Digital Health* of the three prospective controlled pivotal studies (AEYE-1, AEYE-2, AEYE-3, over 1,200 patients) behind its FDA-cleared autonomous AI for diabetic retinopathy, on both handheld and tabletop cameras with one image per eye. Evidence level: this is the right direction, pivotal evidence moved from FDA files into the open literature. Questions before buying: performance in your population and camera fleet versus the trial settings; the gradeability rate on your own operators, since over 99% imageability was achieved under trial conditions; the referral pathway load created by false positives; and what post-deployment monitoring the vendor supports for an autonomous, no-physician-in-the-loop workflow.

Source: [PR Newswire](#)**[PARTNERSHIP] Qscription Technologies and NEOPATHOLOGY sign an MOU to commercialize an FDA-cleared lung imaging AI in the US***Qscription Technologies / NEOPATHOLOGY | 9 September 2026*

What it is: a memorandum of understanding, announced 9 September, to explore joint commercialization and enterprise integration of an unnamed FDA 510(k)-cleared lung imaging AI application. Evidence level: preliminary framework only, no definitive agreement, and the announcement does not identify the cleared device or its K number. The candid premise is worth noting: the companies say cleared AI routinely stalls between clearance and clinical use on IT integration, security review and onboarding. Questions if approached: which device and clearance this concerns, verbatim intended use; who services the integration and who is accountable for post-deployment performance; and what evidence exists beyond the clearance itself.

Source: [PRLog](#)

SECTION 3

POST-MARKET SIGNAL

*Drift | Recalls and MAUDE | Real-world monitoring | EU AI Act Article 72***Vendor-neutral monitoring platforms move into the drift gap the FDA does not cover***Imaging Technology News | 9 September 2026*

An ITN feature on AI sprawl quotes TestDynamics on the governance gap their Satori platform targets: FDA clearance reflects performance evaluated at a point in time, while local populations, scanners, protocols and clinical conditions change, so post-deployment drift is the deployer's problem. The platform integrates, validates and continuously monitors deployed AI tools, triggering alerts when performance decays or when user agreement with AI outputs declines over time. Whatever the tooling choice, the metrics named here, output drift and radiologist-AI agreement trended over time, are exactly the Article 72 evidence stream a deployment file needs; declining reader agreement is often the earliest observable signal.

Link: <https://www.itnonline.com/article/taming-ai-sprawl-radiology>**The post-market evidence gap quantified: three of 1,300+ cleared AI devices evaluated for patient outcomes***TechTarget / PLOS Digital Health | 8 September 2026*

The PLOS Digital Health analysis covered in this week's media block is also a post-market signal: with only 0.9% of FDA-cleared AI devices publishing prospective trial results and 0.2% evaluated for patient-centered outcomes, the burden of demonstrating real-world safety and benefit sits almost entirely on deployers after purchase. Radiology, with 78% of cleared devices and 1% linked to prospective trials, carries the largest share of that unmonitored surface. This is the quantitative case for local concordance monitoring as a standing programme rather than an optional extra.

Link: <https://www.techtarget.com/healthtechnanalytics/news/366650098/Less-than-1-of-FDA...>

SECTION 4

PLAYBOOK SNIPPET

*One concrete step your AI committee can take this month***EU AI Act Article 9 (risk management), pre-deployment**

Step: Before any LLM or classifier pilot, assemble 100 consecutive routine cases from your own service, not teaching cases, and score the candidate tool on them before you see any vendor demo. Record the case list, the model version and the scores in the risk file.

Why: Curated demos systematically overstate: performance on teaching cases does not transfer to routine casemix, and this week's review shows the published literature cannot carry the safety case for you.

Worked example: Wu M et al. show why the local test matters even for well-built tools: ST-USNet, developed on four institutions with validation AUC 0.984 for malignancy, still performed measurably lower on an independent test cohort, and its authors call for prospective multi-center validation before deployment. Academic Radiology, PMID 42705922, DOI 10.1016/j.acra.2026.08.092.

SECTION 5

MEDIA HIGHLIGHTS

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

TechTarget: fewer than 1% of FDA-cleared AI devices have been tested for patient-centered outcomes*TechTarget Healthtech Analytics* | 8 September 2026

Coverage of a PLOS Digital Health study of more than 1,300 FDA-cleared AI-enabled devices: only 34 (2.5%) were linked to registered prospective trials, 12 posted results, and just three devices (0.2%) were evaluated for patient-centered outcomes such as mortality or readmissions. Radiology accounts for 78% of cleared devices but only 1% were linked to prospective trials. The authors propose mandatory trial pre-registration for Class II and III AI devices before clearance. For a deployer the number to remember is the denominator: clearance volume is not an evidence base.

Link: <https://www.techtarget.com/healthtechanalytics/news/366650098/Less-than-1-of-FDA...>**NYU model predicts five-year breast cancer risk from longitudinal 3D mammograms***Medical Xpress / AJR* | 10 September 2026

NYU Langone reported NYU-DRP, a deep learning tool using a woman's serial digital breast tomosynthesis exams to predict five-year breast cancer risk. In the AJR paper it ranked higher-risk women correctly 72% of the time, versus 70% for single-DBT AI, 68% for 2D-mammogram AI and 56% for the Tyrer-Cuzick model on the five-year horizon. Single-center and retrospective, but the direction matters for screening programmes considering risk-adapted intervals: the longitudinal image history, which every screening service already owns, carries risk signal that questionnaire models miss.

Link: <https://medicalxpress.com/news/2026-09-ai-imaging-tool-breast-cancer.html>**MedCity News: radiology at a structural breaking point, and why bolt-on AI will not fix it***MedCity News* | 8 September 2026

An operations-focused essay arguing that layering point AI solutions onto fragmented radiology IT delivers limited gains, and that the coming PowerScribe 360 sunset forces a bigger question: not which reporting tool to buy, but what interpretation environment to build. It cites The Imaging Wire's 2026 trend list (AI workflow commoditization, bundling into care-pathway solutions, accelerating M&A) and a published evaluation finding AI gains in accuracy and standardization constrained by generalizability, interpretability and organizational readiness. Useful framing for any 2027 platform decision.

Link: <https://medcitynews.com/2026/09/is-radiology-at-a-structural-breaking-point-what...>

SECTION 6

PROFESSIONAL SOCIETIES

SIIM | ACR | RCR | Official publications and event coverage

Neiman Institute and JACR publish real-world shadow-mode evaluation of aneurysm AI across a health system

8 September 2026

The Harvey L. Neiman Health Policy Institute announced a JACR study from Northwell Health this week: a prospective shadow-mode evaluation of an FDA-cleared intracranial aneurysm algorithm on 3,856 CTA examinations. AI and radiologists agreed in over 96% of cases; AI added 55 true-positive aneurysms radiologists had not reported (a 39% relative detection increase) at the cost of 46 false positives, while radiologists caught 30 aneurysms the AI missed. Performance varied sharply by care setting, favorable inpatient, marginal outpatient. The authors' conclusion is the governance lesson: evaluate AI on how it changes physician performance in your own settings, and keep monitoring after implementation.

Source: <https://www.newswise.com/articles/ai-helps-radiologists-detect-39-more-brain-aneurysm-cases-in-real-world-study>

RSNA Board advances advanced practice provider integration as a workforce strategy

8 September 2026

RSNA published a Board-level update on 8 September on integrating advanced practice providers (nurse practitioners, physician assistants, radiologist assistants) into radiologist-led care teams, following its APP task force report accepted in June 2025. The Board is prioritizing convening and education, with APP-targeted programming in development for the annual meeting, while deferring standards-setting and advocacy given differing scopes of practice across states. Relevant to AI governance planning because the same capacity pressure driving APP integration is the stated business case for most imaging AI purchases.

Source: <https://www.rsna.org/news/2026/september/advanced-practice-providers>

SECTION 7

KEY OPINION LEADERS

LinkedIn | Public statements | Week of 6 to 12 September 2026

Bernardo Bizzo

Associate Chief Science Officer, ACR Data Science Institute; MD-PhD radiologist

Posted 9 September (amplified by Woojin Kim, CMO of ACR DSI): 'The first wave of imaging AI produced a flag. The next wave will produce a draft report. Monitoring has to keep up.' Announcing a new JACR paper from the ACR DSI team, 'AI Monitoring AI with LLMs: The American College of Radiology's Imaging AI Registry', describing the LLM engine inside Assess-AI that extracts findings from radiology reports as the continuous monitoring signal for deployed imaging AI, spanning dozens of use cases and hundreds of thousands of report-AI pairs. His point on why now: as agentic workflows start generating draft reports that radiologists review and sign, the monitoring question changes.

Source: [LinkedIn](#)

Amine Korchi

Radiologist; HealthTech founder and investor; Geneva

Posted 7 September on Vara's CE certification (Class IIb) for autonomous mammography AI, the first autonomous AI in mammography and the second in radiology after Oxipit's chest x-ray clearance. He calls it a landmark but flags the deployer's problem precisely: he could not find Vara's official intended use and CE documentation, and judges from the announcement that it conceptually mirrors Oxipit, AI autonomously clears normal exams while abnormal or uncertain cases go to a radiologist, with continuous monitoring layered on top. A working demonstration of the first question to ask any autonomous AI vendor: show me the certified intended-use text.

Source: [LinkedIn](#)

Louis Blankemeier

Co-Founder and CEO, Cognita; radiology AI researcher

Posted 8 September (amplified by Nina Kottler, Chief Medical AI Officer, Mosaic Clinical Technologies): a core-thesis post arguing human intelligence and artificial intelligence are fundamentally different designs whose strengths are non-overlapping, so expecting one to be better at everything defies logic; in practice AI is impressive on many tasks but makes mistakes humans do not make, which is the design argument for radiologist-AI complementarity rather than replacement. Kottler's repost signals how practice-level medical AI leadership is framing the division of labor.

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 6 to 12 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.