

ЕЖЕНЕДЕЛЬНЫЙ ОБЗОР ИИ

Радиология и медицинская визуализация

Неделя 32 | 2 по 8 августа 2026

Д-р Сергей Морозов | drsergeymorozov.com

Составлено по публично доступным источникам (индексируемая рецензируемая литература, официальные пресс-релизы, регуляторные базы данных). Не является медицинской, регуляторной или финансовой консультацией.

8

Рецензируемые статьи

4

Индустрия / FDA

3

Обзор СМИ

3

Лидеры мнений

SECTION 0

КРАТКОЕ РЕЗЮМЕ

Ключевые тренды, неделя с 2 по 8 августа 2026

[БЕНЧМАРК LLM] МУЛЬТИМОДАЛЬНЫЕ LLM ЗАСТРЯЛИ НИЖЕ КЛИНИЧЕСКОЙ ПРИГОДНОСТИ В РАДИОЛОГИИ

Wu Q и соавт. (Journal of Medical Internet Research) построили RadM-Bench - 720 двуязычных случаев по 9 субспециальностям - и оценили 10 мультимодальных LLM по шкале от 0 до 3. Все остались ниже 1,5. Ключевые 2D-изображения, отобранные рентгенологом, помогли каждой модели (+19,8% - +139,2%), но объемные данные ухудшили все оцениваемые модели, а результаты на учебных случаях не перенеслись на рутинные китайские клинические случаи (ОЗ с изображениями: с 1,14 до 0,79). Перевод клинических анамнезов на английский ухудшил все 10 моделей.

[КЛИНИЧЕСКИЙ ПРОГНОЗ] ТРАНСФУЗИОННАЯ МОДЕЛЬ, ЗНАЮЩАЯ СОБСТВЕННУЮ НЕОПРЕДЕЛЕННОСТЬ

Li X и соавт. (JMIR Medical Informatics) разработали MCGB - градиентный бустинг с клиническими ограничениями, предсказывающий потребность в гемотрансфузии и ее дозу у 849 пациентов с кровотечением из верхних отделов ЖКТ в 3 больницах. AUROC 0,97, чувствительность 0,99, специфичность 0,87 на полностью исключенной из разработки больнице; дозы сопровождаются конформными 95% интервалами предсказания (покрытие 0,94). Монотонные ограничения и количественная неопределенность закрывают два стандартных возражения против ML у постели пациента.

[СМЕЩЕНИЯ] РАЗНООБРАЗИЕ ДАННЫХ ВАЖНЕЕ ИХ ОБЪЕМА В СЕГМЕНТАЦИИ ГЛИОБЛАСТОМЫ

Lee RS и соавт. (AJNR) провели аудит четырех моделей nnUNet для сегментации глиобластомы на 480 независимых пациентах. Модель, обученная только на белых мужчинах неиспаноязычного происхождения, показала худшие результаты (Dice 0,943 FLAIR, 0,909 контрастное усиление) и смещения по возрасту и курению; умеренная по размеру, но разнообразная модель BraTS 2024 показала лучшие (0,996 и 0,999) и обошла значительно более крупную FeTS. Гетерогенность снижает смещение даже при одинаковом объеме обучающей выборки.

[ОПЕРАЦИИ] ЯЗЫКОВЫЕ МОДЕЛИ НАХОДЯТ ДЕНЬГИ, УТЕКАЮЩИЕ ИЗ УЗИ-ЗАПИСЕЙ

Nguyen K и соавт. (Journal of Medical Internet Research) с помощью BioClinBERT (точность 0,97) выявляли процедуры УЗИ в месте оказания помощи в 559 029 акушерско-гинекологических обращениях, а затем показали, что стандартизированный шаблон документирования снизил долю пропущенных счетов с 10,0% до 2,4% (ОШ 0,22, P<.001), причем 75,4% кодов формировались шаблонами автоматически. ИИ в роли аудитора выручки - с проверкой, что выигрыш дал рабочий процесс, а не рост числа процедур.

[ФОКУС НЕДЕЛИ] ГЛАВНОЕ ПРОИСХОДИТ НА РАБОЧЕЙ СТАНЦИИ, А НЕ В МОДЕЛИ

Прочитанные вместе, статьи недели смещают фронт от точности к внедрению. Wu Q и соавт. показывают, что мультимодальные LLM остаются ниже клинически пригодного уровня, как только случаи перестают быть отобранными. Zhang R и соавт. находят, что стратегия промптинга должна подбираться под архитектуру модели, а режимы рассуждения и число параметров могут отнимать качество при классификации описаний. Moturu A и соавт. наблюдают, как два рентгенолога работают с подсветкой на детской МРТ всего тела: инструмент замедляет обоих и меняет их разметку в противоположных направлениях, откуда вывод о возможной необходимости калибровки ИИ под каждого читателя. Даже самый сильный клинический результат недели - трансфузионная модель MCGb Li X и соавт. - обязан доверием ограничениям, калибровке и межцентровому тестированию, а не одной лишь дискриминации. Закупочный вопрос недели - не сколько баллов набирает модель, а что происходит с людьми, рабочим процессом и счетами, когда она приходит.

SECTION 1

РЕЦЕНЗИРУЕМЫЕ СТАТЬИ

Источник: PubMed | Индексация проверена | Без препринтов | 2 по 8 августа 2026

Prehospital Injury Severity Estimate (PHISE) matches in-hospital trauma scores when embedded in AI models.

Sigle M, et al. | *NPJ digital medicine* | 2026-08-07

PHISE - догоспитальная шкала тяжести травмы, построенная путем проекции описаний повреждений с кодами AIS из американского National Trauma Data Bank на восемь анатомических зон и четыре уровня тяжести; инференс с помощью LLM использовался, чтобы определить, какие повреждения требуют визуализации, а какие можно оценить клинически. Цель - анатомический компонент тяжести, применимый до любой визуализации и изначально спроектированный как входная переменная для моделей машинного обучения вместе с другими догоспитальными данными.

Ключевые показатели: Абстракт приводит сравнения качественно, без числовых значений: в когорте NTDB шкала PHISE показала сильную корреляцию и хорошую калибровку с ISS и NISS, но более низкую предсказательную способность для госпитальной летальности и потребности в гемотрансфузии. Во внешней валидации на многонациональном регистре TraumaRegister DGU она уступила ISS/NISS по летальности, сравнялась по прогнозу гемотрансфузии, а при включении PHISE в ML-модели с догоспитальными переменными разрыв сократился еще больше.

Клиническая значимость: Решения о сортировке при травме принимаются до визуализации, тогда как все эталонные шкалы от нее зависят. Анатомическая шкала, независимая от визуализации, изначально задуманная как вход для моделей ИИ и прошедшая внешнюю валидацию на регистре другой травматологической системы, - реальный путь к более ранней стратификации и более разумной активации ресурсов визуализации и службы крови.

PMID 42562846 DOI 10.1038/s41746-026-03074-7

Prediction of Blood Transfusion Need and Dose in Patients With Upper Gastrointestinal Bleeding: Retrospective Multicenter Prediction Model Study.

Li X, et al. | *JMIR medical informatics* | 2026-08-04

Двухэтапный градиентный бустинг с клиническими ограничениями (MCGb), который сначала предсказывает, потребуется ли пациенту с эндоскопически подтвержденным кровотечением из верхних отделов ЖКТ гемотрансфузия, а затем оценивает дозу с конформными 95% интервалами предсказания. Разработан на ретроспективной многоцентровой когорте из 849 взрослых из 3 больниц Чунцина (2019-2025) с монотонными клиническими ограничениями, тестированием на полностью исключенной из разработки больнице и прототипом интерфейса для использования у постели пациента.

Ключевые показатели: AUROC 0,97 и AUPRC 0,91 для потребности в гемотрансфузии; при пороге 0,50 чувствительность 0,99, специфичность 0,87, F1 0,85. Предсказание дозы у трансфузированных пациентов: R2 0,95, средняя абсолютная ошибка 0,04, покрытие 95% интервалов предсказания 0,94. Оценка проводилась на независимой тестовой больнице, с попеременным внешним тестированием между больницами как проверкой устойчивости.

Клиническая значимость: Пороги гемотрансфузии при кровотечениях из верхних отделов ЖКТ остаются спорными, а фиксированные пороги гемоглобина игнорируют конкретного пациента. Калиброванная модель, которая количественно оценивает собственную неопределенность, соблюдает клиническую монотонность и протестирована между центрами, отвечает на два главных упрека к ML-поддержке решений - непрозрачность и самоуверенность - и имеет прямое значение как для решений у постели пациента, так и для планирования службы крови.

PMID 42550984 DOI 10.2196/83889

Evaluating Sociodemographic Biases in Artificial Intelligence-Based Glioblastoma Response Assessment Algorithms.

Lee RS, et al. | *AJNR. American journal of neuroradiology* | 2026-08-03

Аудит смещений четырех моделей сегментации глиобластомы на MPT на базе nnUNet, обучающие выборки которых различаются размером и демографическим составом: послеоперационная модель FeTS (более 10 000 исследований, многонациональная), модель BraTS 2024 (более 2 000 исследований, Северная Америка) и две малые одноцентровые модели (более 200 исследований каждая) - демографически однородная и гетерогенная. Все оценивались на независимом, вручную откорректированном наборе из 480 проспективно набранных пациентов.

Ключевые показатели: Модель, обученная исключительно на белых мужчинах неиспаноязычного происхождения, показала самые низкие коэффициенты Dice (0,943 для FLAIR, 0,909 для накапливающей контраст опухоли) и смещения по возрасту и статусу курения. Модель BraTS показала самые высокие Dice (0,996 FLAIR, 0,999 контрастное усиление) и наименьшее смещение, превзойдя модель FeTS несмотря на то, что у FeTS самая большая обучающая выборка. Социодемографические эффекты анализировались бета-регрессией.

Клиническая значимость: Для тех, кто собирает обучающие данные, важны два результата: демографическая гетерогенность снижает смещение даже без увеличения объема выборки, а умеренная по размеру, но разнообразная когорта побеждает значительно большую. Смещение моделей сегментации в целом невелико, но концентрируется там, где обучающие данные однородны, - аргумент в пользу аудита состава, а не только объема данных при закупках и валидации.

PMID 41663205 DOI 10.3174/ajnr.A9217

Machine Learning to Identify Point-of-Care Ultrasound and Evaluate Standardized Documentation: Retrospective Operational Cohort Study.

Nguyen K, et al. | *Journal of medical Internet research* | 2026-08-07

Операционное исследование в крупном академическом центре: ML (LightGBM и BioClinBERT) использовали для поиска процедур ультразвука в месте оказания помощи (POCUS), скрытых в свободном тексте акушерско-гинекологических записей, а затем измерили эффект стандартизированного шаблона документирования процедур (ProcDoc) на полноту выставления счетов. Анализ охватил 559 029 обращений 109 776 пациенток в 11 клиниках с 2018 по 2024 год; шаблон внедрен в феврале 2023 года.

Ключевые показатели: BioClinBERT достиг точности 0,97 (F1 0,55-0,63) в выявлении задокументированных и пропущенных процедур. Внедрение ProcDoc достигло 75,1% за 12 месяцев. Доля счетов, пропущенных врачами и выявленных постфактум, снизилась с 10,0% до 2,4% (отношение шансов 0,22, 95% ДИ 0,17-0,30, P<0,001), при общем росте выставления счетов за POCUS на 0,6%. 1 812 из 2 404 кодов CPT после внедрения (75,4%) сформированы шаблонами автоматически.

Клиническая значимость: Здесь ИИ выступает инструментом аудита, а не диагностики: модели количественно оценили потерю выручки из-за незадокументированных исследований POCUS и подтвердили, что выигрыш дала перестройка рабочего процесса, а не рост объема процедур. Для руководителей службы лучевой диагностики переносимый урок - именно связка: языковая модель для измерения проблемы в свободном тексте плюс структурированное документирование для ее устранения.

PMID 42566781 DOI 10.2196/84454

A Bilingual Benchmark for Evaluating Diagnostic Performance of Multimodal Large Language Models in Radiology (RadM-Bench): Evaluation Development and Validation.

Wu Q, et al. | *Journal of medical Internet research* | 2026-08-07

RadM-Bench - двуязычный радиологический бенчмарк из 720 случаев по 9 субспециальностям: 360 англоязычных публичных учебных случаев с преобладанием редких болезней и 360 китайских рутинных клинических случаев. Десять мультимодальных LLM (4 проприетарные, включая GPT-4o, O3 и варианты Gemini; 6 открытых) тестировались с анамнезом без изображений, с ключевыми 2D-изображениями, отобранными рентгенологом, и с объемными данными при 2 и 10 кадрах в секунду; ответы оценивали от 0 до 3 два сертифицированных рентгенолога вслепую.

Ключевые показатели: Средние оценки всех 10 моделей остались ниже 1,5 из 3 на обоих наборах. Добавление ключевых 2D-изображений улучшило каждую модель (+19,8% - +139,2%), но объемный ввод при 10 кадрах/с ухудшил все 8 оцениваемых моделей относительно 2D-базы (-5,3% - -31,4%). Проприетарные модели теряли в качестве при переходе от учебных к клиническим случаям (O3 с изображениями: 1,14 против 0,79), тогда как китайско-ориентированные открытые модели улучшались. Перевод китайских анамнезов на английский снизил оценки всех 10 моделей, статистически значимо у 9.

Клиническая значимость: Бенчмарк выявляет именно те разрывы, с которыми столкнулось бы клиническое внедрение: модели пока не умеют работать с 3D-данными, результаты на отобранном учебном материале превышают ожидания для рутинных, а преимущество на редких болезнях, видимое на учебных наборах, в клинике обращается вспять. Языковой результат опровергает допущение, что ввод на английском всегда оптимален. Полезный ориентир перед любым пилотом мультимодальной LLM в описании исследований или триаже.

PMID 42566748 DOI 10.2196/92183

Large Language Models for Classifying Usual Interstitial Pneumonia from Radiology Reports: Native Reasoning Versus Structured Prompting.

Zhang R, et al. | *Journal of imaging informatics in medicine* | 2026-08-05

Систематическое сравнение 10 открытых LLM семейств Llama, Qwen и Gemma (от 8 до 405 млрд параметров) для классификации паттернов обычной интерстициальной пневмонии (ОИП) по 270 описаниям КТ высокого разрешения согласно критериям Fleischner; референсом служил консенсус двух старших торакальных рентгенологов. Три стратегии промптинга скрещивались с архитектурой модели, четыре модели с нативным рассуждением дополнительно тестировались в режиме размышления.

Ключевые показатели: Лучшая конфигурация достигла каппы Козна 0,70 и точности 82% в четырехклассовой задаче. Структурированный промптинг с рассуждением улучшил все модели Llama, но ухудшил все модели с нативным рассуждением (Qwen 3.5, Gemma 4). Режим размышления вредил промптам на основе критериев и ни разу не дал лучшей конфигурации. Llama 3.1 405B не дала преимущества перед Llama 3.3 70B, а Qwen на 397 млрд параметров уступила плотной модели на 27 млрд.

Клиническая значимость: Для конвейеров разметки описаний операционные выводы таковы: промпт должен подбираться под архитектуру, режимы рассуждения могут активно вредить классификации по критериям, а число параметров - плохой критерий закупки. Команды, размечающие описания для обучения моделей, могут получить уровень state-of-the-art на открытых моделях среднего размера, если тестируют пару промпт-архитектура, а не полагаются на то, что новее и больше значит лучше.

PMID 42557432 DOI 10.1007/s10278-026-02178-6

Automatic Patient Eligibility for Photon-counting CT Using Discriminative and Generative AI Models in Neuroradiology.

Martin-Noguerol T, et al. | *Academic radiology* | 2026-08-05

Проверка концепции NLP-маршрутизатора, который читает испаноязычные направления на КТ в нейрорадиологии и решает, требует ли исследование полного протокола фотонно-счетной КТ (PCCT) или может быть выполнено на обычном томографе с энергоинтегрирующим детектором. 800 направлений (2012-2025) разметили два нейрорадиолога (каппа Козна 0,74); дообученные дискриминативные трансформеры сравнивались с генеративными моделями в режиме zero-shot (ChatGPT 5.2, Gemini 3) при пятикратной перекрестной валидации.

Ключевые показатели: Лучшей оказалась биомедицинская RoBERTa: точность 0,925, precision 0,906, полнота 0,967, F1 0,935, AUC 0,920 - без значимого отличия от испанской BERT, но со статистически значимым превосходством над Llama-3.1-8B, Mistral-7B и обеими генеративными LLM, показавшими худшие результаты.

Клиническая значимость: Мощности PCCT дефицитны, а руководств о том, когда ее предпочесть, нет, поэтому автоматический маршрутизатор решает реальную задачу распределения ресурсов. Результат несет и более широкий сигнал для ИТ лучевой диагностики: малые дообученные дискриминативные модели по-прежнему превосходят универсальные генеративные LLM в четко очерченной классификации - при доле их стоимости и регуляторной нагрузки.

PMID 42557204 DOI 10.1016/j.acra.2026.07.026

Human-AI Interaction With AI-Assisted Tumor Overlays in Pediatric Whole-Body Magnetic Resonance Imaging: Exploratory Reader Study.

Moturu A, et al. | *JMIR human factors* | 2026-08-07

Пилотное исследование взаимодействия человека и ИИ: патч-ориентированная сегментационная подсветка выделяет опухолеподобные зоны на скрининговой МРТ всего тела у детей с синдромом Ли-Фраумени; модель обучена на аугментированных 2D-срезах из 675 объемов МРТ всего тела. Два рентгенолога читали исследования с подсветкой и без нее; фиксировались время оценки, отметки кандидатных очагов, рекомендации по дообследованию и субъективная нагрузка.

Ключевые показатели: Абстракт приводит результаты только качественно, без числовых значений: среднее время оценки исследования с инструментом выросло у обоих читателей; один стал отмечать больше кандидатных очагов, другой - меньше; один сообщил о снижении умственной нагрузки, оба - о снижении стресса при работе с инструментом. Межэкспертная вариабельность была выраженной на всем протяжении исследования.

Клиническая значимость: Честный портрет того, что ИИ-ассистент реально делает на рабочей станции: он изменил поведение читателей в противоположных направлениях и замедлил обоих, одновременно снизив стресс. Вывод авторов о том, что инструменты ИИ, возможно, требуют индивидуальной калибровки под каждого рентгенолога, ставит под сомнение модель развертывания по принципу один инструмент для всех, а детский скрининг редких синдромов - именно та область, где цена второй пары глаз максимальна.

PMID 42567688 DOI 10.2196/81066

SECTION 2

ИНДУСТРИЯ И РЕГУЛИРОВАНИЕ

Источники: Официальные реестры | Пресс-релизы | Базы данных FDA

[MARKET] GE HealthCare выпускает Invenia ABUS Prime и одобренный FDA виевер ABUS StreamVue

GE HealthCare | 3 августа 2026

GE HealthCare объявила о запуске Invenia ABUS Prime - новой системы автоматизированного УЗИ молочных желез для дополнительного скрининга при плотной ткани - вместе с ABUS StreamVue, браузерным корпоративным виевером без локальной установки, недавно получившим разрешение FDA 510(k) и маркировку CE. ABUS Prime включает ИИ-оценку качества сканирования, автоматическое определение соска и поддержку чтения QVCAD; по материалам компании, ИИ-ассистированный просмотр достигает чувствительности 93% и сокращает время интерпретации до 33%, а Fast Scan уменьшает время сканирования до 40%. StreamVue, созданный с Candelis, рассчитан на удаленное чтение и стандартизацию между площадками; одна из первых инсталляций Invenia ABUS Prime в США - в центре BeSound в Лос-Анджелесе.

Source: [GE HealthCare](#)

[REGULATION] Cortechs.ai расширяет маркировку CE по EU MDR на всю платформу количественной визуализации

Cortechs.ai | 4 августа 2026

Cortechs.ai объявила о расширении европейской маркировки CE по регламенту EU MDR на новейшие продукты: версию 5 NeuroQuant (включая NeuroQuant MS и NeuroQuant ARIA), NeuroQuant Brain Tumor и OnQ Prostate. Сертификация открывает коммерциализацию текущего поколения инструментов автоматической волюметрии мозга, мониторинга опухолей мозга и визуализации простаты на основе RSI на всей территории Европейской экономической зоны. Компания называет этот шаг фундаментом европейской стратегии выхода на рынок после более чем года регуляторной работы по MDR.

Source: [PR Newswire](#)

[PARTNERSHIP] 4DMedical инвестирует 3,4 млн долларов в RevealDx и получает глобальную дистрибуцию RevealAI-Lung

4DMedical | 3 августа 2026

4DMedical объявила о стратегической инвестиции 3,4 млн долларов в RevealDx (Сизтл) и дистрибьюторском соглашении для США, Европы, Австралии и Новой Зеландии по RevealAI-Lung - CADx-инструменту, характеризующему легочные узлы с помощью индекса Malignancy Similarity Index (mSI). RevealAI-Lung имеет разрешение FDA, сертификацию EU MDR и одобрение TGA, валидирован более чем на 1 500 пациентах и возмещается в США по кодам CPT 0721T и 0722T. 4DMedical встроит его в недавно приобретенную платформу КТ грудной клетки contextflow и объединит со своей функциональной визуализацией легких CT:VQ, собирая единую кардиопульмональную экосистему: выявление, стратификация риска и функция.

Source: 4DMedical

[MARKET] GE HealthCare представляет семейство компактных УЗИ-систем LOGIQ e с функциями ИИ

GE HealthCare | 4 августа 2026

GE HealthCare представила LOGIQ e Xi и LOGIQ e Si - ультразвуковые системы в формате ноутбука, сочетающие производительность консольных аппаратов с ИИ-функциями для мультидисциплинарной работы и диагностики в месте оказания помощи: разметка анатомических ориентиров Nerveblox для блокад нервов, распознавание структур MSK Go, Automated Function Imaging для оценки сердца и фетальная биометрия SonoBiometry. Инструменты Scan Assistant и Voice Control сокращают ручные операции; Xi поддерживает облачное управление парком Verisound и беспроводные датчики Vscan Air. Запуск расширяет сегмент портативной визуализации с ИИ в неотложной помощи, интенсивной терапии, МСК и анестезиологии.

Source: MPO Magazine

SECTION 3

ОБЗОР СМИ

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

Могут ли LLM закрыть разрыв в коммуникации с направляющими врачами?

AuntMinnie.com | 4 августа 2026

Обзор исследования UCSF, опубликованного 4 августа в Radiology: LLM способны формировать более полные и клинически полезные показания к визуализации, чем указывают направляющие врачи. На материале 28 313 пациентов, 740 867 клинических записей и 77 626 исследований команда сравнила проприетарные и открытые модели; в читательском исследовании с 20 рентгенологами Claude 3.5 Sonnet получил высшую оценку полноты в 37,14% случаев и фактической точности в 68,05% против 6,64% и 50% у направляющих врачей; открытая Qwen 2.5-7B шла следом. Сопровождающая редакционная статья подчеркивает прозрачность, мониторинг и контроль человека до внедрения.

Ссылка: <https://www.auntminnie.com/imaging-informatics/artificial-intelligence/article/1...>

Konica Minolta расширяет платформу Eха Enterprise облачной телерадиологией и ИИ-партнерами

Radiology Business | 3 августа 2026

Материал о расширении Konica Minolta Medical Imaging USA своего PACS Eха Enterprise до облачной телерадиологической платформы, собранной из партнерских решений, а не собственных разработок: оркестрация потоков NewVue, генеративные ИИ-описания RADPAIR, облачная инфраструктура AWS Health Imaging и VNA Apollo. Браузерная платформа тарифицируется за исследование и масштабируется от практики одного врача до больничных систем на 500оек - иллюстрация тренда платформизации, при котором вендоры PACS интегрируют сторонний ИИ вместо собственной разработки.

Ссылка: <https://radiologybusiness.com/topics/health-it/enterprise-imaging/konica-minolta...>

Итоги недели: сроки подготовки заключений продолжают расти

AuntMinnie.com | 7 августа 2026

Еженедельная редакционная заметка AuntMinnie выделяет данные Harvey L. Neiman Health Policy Institute: сроки выдачи радиологических заключений для пациентов Medicare продолжили расти в 2024 году, хотя и медленнее, чем в 2023-м, на фоне дефицита рентгенологов и растущих объемов. Также отмечен пилот UMass Memorial Health, направлявший неотложных пациентов первичного звена сразу в отделение лучевой диагностики вместо приемного отделения, что сократило время от прибытия до результатов визуализации почти на 50% - организационный результат, напрямую связанный с аргументом о пропускной способности, движущим внедрение ИИ в радиологии.

Ссылка: <https://www.auntminnie.com/editors-note/article/15831989/week-in-review-imaging-...>

SECTION 4

ПРОФЕССИОНАЛЬНЫЕ ОБЩЕСТВА

SIIM | ACR | RCR | Официальные публикации и освещение событий

RSNA запускает Knee Abnormality Detection AI Challenge 2026

5 августа 2026

RSNA выпустила официальный пресс-релиз своего AI Challenge 2026 - первого, посвященного МРТ опорно-двигательного аппарата, и первого, объединяющего изображения с текстом радиологических заключений. Обучающий набор включает более 5 000 МРТ коленного сустава с заключениями на 12 языках из 16 центров на пяти континентах; оценочный набор аннотирован экспертами-рентгенологами. Соревнование на платформе Kaggle продлится до 22 октября; призовой фонд 77 000 долларов, впервые с наградой за самые эффективные модели; победителей отметят на конгрессе RSNA 2026 в Чикаго (29 ноября - 3 декабря).

Source: <https://www.rsna.org/media/press/2026/2669>

SECTION 5

ЛИДЕРЫ МНЕНИЙ

LinkedIn | Публичные высказывания | Неделя с 2 по 8 августа 2026

Curt Langlotz

Профессор радиологии, медицины и биомедицинских данных, Stanford University School of Medicine

7 августа поделился публикацией MedVAL в npj Digital Medicine - статьи Стэнфорда, соавтором которой он является: Toward expert-level medical text validation with language models, опубликована 4 августа. Работа посвящена практическому узкому месту генеративного ИИ в радиологии - проверке фактического соответствия сгенерированного медицинского текста исходным данным - и предлагает сами языковые модели в роли масштабируемых валидаторов. На этой неделе Langlotz также написал о встрече с командой Raidium в Пало-Альто, вновь процитировав свою формулу 2017 года: ИИ не заменит рентгенологов, но рентгенологи, использующие ИИ, заменят тех, кто его не использует.

Source: [LinkedIn](#)

Woojin Kim

Chief Strategy Officer и CMIO компании HOPPR, CMO ACR Data Science Institute, МСК-рентгенолог

6 августа сообщил, что примет участие в EuSoMI 2026 в Ираклионе на Крите (9-10 октября), представляя ACR Data Science Institute, и поделился кампанией института о его вкладе в будущее ИИ в радиологии. Ранее на неделе (2 августа) рассказал о своей лекции о предпринимательстве и ИИ в радиологии в курсе Healthcare Innovation and Transformation Джеральдин МакГинти в Weill Cornell - знак того, что лидеры радиологического ИИ входят в программы обучения руководителей здравоохранения.

Source: [LinkedIn](#)

Amine Korchi

Рентгенолог и healthtech-предприниматель, Женева; соавтор обзора VLM нейрорадиологической команды HUG

7 августа написал о только что опубликованном систематическом обзоре, подготовленном с командой нейрорадиологии Университетских госпиталей Женевы (HUG), о базовых моделях vision-language для интерпретации 3D-нейровизуализации. Его сводка доказательств: специализированные нейровизуализационные VLM устойчиво превосходят универсальные мультимодальные модели и приближают ИИ-ассистированное составление заключений к клинической реальности, но к рутинному внедрению пока не готовы. Среди названных пробелов: слишком много правок заключений для реального выигрыша во времени, огромные вычислительные затраты, отсутствие клинически значимых метрик оценки, галлюцинации и пропущенные находки, интеграция в рабочий процесс и регуляторная зрелость. Для разработчиков ИИ он представляет эти ограничения как список задач, на котором будет построено следующее поколение компаний радиологического ИИ.

Source: [LinkedIn](#)

ПРИМЕЧАНИЕ О КОНТРОЛЕ КАЧЕСТВА

Подготовлено д-ром Сергеем Морозовым по публично доступным источникам: рецензируемые статьи, индексируемые в PubMed (PMID и DOI проверены, без препринтов и arXiv), официальные пресс-релизы, регуляторные базы данных и профильные СМИ. Не является медицинской, регуляторной или финансовой консультацией. Даты публикации всех статей проверены и попадают в период с 2 по 8 августа 2026. Отбор статей детерминирован (фиксированный запрос в PubMed, белый список журналов Q1 и флагманских изданий, исключение чистых обзоров [метаанализы сохраняются], работ только по радиомике и по интервенционной радиологии, затем прозрачная оценка ранжирования). Отраслевые новости взяты из официальных пресс-релизов и регуляторных уведомлений; материалы лидеров мнений из публичной активности в LinkedIn в пределах указанного периода.