

Готова ли ваша организация к управлению ИИ в сфере нормативно-правового регулирования?

Источник: MedTech Intelligence

Оригинал: https://medtechintelligence.com/feature_article/is-your-organization-ready-to-govern-ai-in-regulatory-affairs/

медтех

нормативно-правовое регулирование

регулирование

риски

управление ИИ

Внедрение — это не управление

Искусственный интеллект уже внедрен в функции по регуляторным вопросам (Regulatory Affairs, RA) во всей индустрии медицинских технологий (medtech). Команды используют его для обобщения руководящих документов, поддержки подготовки черновиков на первом этапе, сравнения версий маркировки и сортировки поступающей аналитической информации. Во многих организациях это происходило незаметно: сотрудники находили инструменты, которые помогали им работать быстрее, и темпы внедрения опережали политики, призванные ими управлять. **«Теневой ИИ» (Shadow AI)** — это практическое название того, как это выглядит на местах. Член команды, находясь под давлением дедлайна, открывает публичный ИИ-инструмент и составляет проект ответа на замечания регулятора. Результат выглядит профессионально и напоминает текст, написанный специалистом по регуляторным вопросам. Он проходит проверку, при этом никто не отмечает ни используемый инструмент, ни промпт (запрос), ни версию, ни этап верификации. В этот раз всё прошло гладко. Но нет никаких записей, которые позволили бы организации объяснить, что произошло бы, если бы что-то пошло не так.

Регуляторные вопросы — это не та функция, где достаточно результата, который просто «выглядит правдоподобно». Когда ИИ влияет на подачу документов, ответ на замечания, обсуждение маркировки или оценку изменений, должны быть даны ответы на три вопроса: какой инструмент использовался? Как был проверен результат? Кто нес за это ответственность? В настоящее время большинство организаций не могут последовательно ответить на все три вопроса. Вопрос не в том, может ли ИИ ускорить регуляторную работу, так как большинство опытных профессионалов согласятся, что может. Вопрос в том, разработала ли организация систему управления (**governance**), соответствующую темпам внедрения.

Практический способ оценки вашего текущего положения

Операционная модель зрелости возможностей ИИ для регуляторных вопросов предлагает пятиуровневую структуру: **Неструктурированная (Unstructured)**, **Определенная (Defined)**, **Контролируемая (Controlled)**, **Управляемая (Governed)** и **Оптимизированная (Optimized)**, которую любая medtech-организация может использовать для оценки текущей практики и определения того, что необходимо развивать далее, как показано на Рисунке 1. Эти уровни — не целевые ориентиры, а диагностический инструмент. Цель состоит не в том, чтобы достичь уровня «Оптимизированная». Цель — точно знать, где организация находится сегодня, и как выглядит наиболее значимый следующий шаг.

Рисунок 1: Модель зрелости возможностей регуляторного ИИ

Таблица 1 обобщает, как каждый уровень выглядит на практике и в функции RA medtech-компании. Важно помнить, что эти уровни описывают организацию, а не инструмент. Команда, использующая сложную ИИ-платформу без документированных сценариев использования и без стандартов проверки, всё равно остается на **Неструктурированном** уровне. Возможности инструмента не тождественны зрелости организации, которая его использует.

Таблица 1: Операционная модель зрелости возможностей ИИ для регуляторных вопросов

Уровень	Что это значит	Как это выглядит в функции RA medtech-компании	Основной риск
Неструктурированная	Использование ИИ носит неформальный характер и незаметно для организации. Нет утвержденных сценариев использования, правил обработки данных, ожиданий от проверки, нет различия между корпоративными и публичными инструментами.	Специалист использует публичную LLM (большую языковую модель) для подготовки ответа на вопрос компетентного органа или обобщения руководства MDCG. Не существует правил относительно того, что можно вводить, как должен проверяться результат или сохраняются ли данные.	Фактические ошибки, утечка данных и отсутствие защищенного аудиторского следа проверки
Определенная	Выявлены начальные сценарии использования и начинают формироваться базовые рекомендации; однако практика неоднородна, и риски еще не дифференцированы по задачам.	Команды комфортно используют ИИ для резюмирования встреч и сортировки литературы, но тот же неформальный подход распространяется на подготовку черновиков	Высокорисковые задачи выполняются с той же неформальностью, что и низкорисковые

Уровень	Что это значит	Как это выглядит в функции RA medtech-компании	Основной риск
		ответов на замечания — задач, которые влекут за собой совершенно разные последствия для комплаенса.	
Контролируемая	Документированы разрешенные и запрещенные способы использования. Внедрено обучение на основе ролей. Ожидания от проверки стандартизированы и больше не зависят от индивидуального суждения.	Сотрудники понимают, что обобщение публичного руководства не требует такой же строгости проверки, как подготовка ответа на запрос FDA по замечаниям. Для регулируемого контента требуются одобренные корпоративные инструменты; формируются стандарты документации.	Результаты ИИ все еще могут не отслеживаться последовательно, когда они влияют на регулируемый контент
Управляемая	Использование ИИ встроено в регуляторную операционную модель. Ответственность четко распределена,	Рабочие процессы с поддержкой ИИ для подачи документов, проверки маркировки и переписки с	Управление существует на бумаге, но может еще не отражать то, как инструменты используются в

Уровень	Что это значит	Как это выглядит в функции RA medtech-компании	Основной риск
	определены пути эскалации, а утверждение сценариев использования привязано к уровню риска. Налажено кросс-функциональное взаимодействие с отделами качества, IT, юридическим отделом и отделом защиты данных.	регуляторами официально утверждены. Организация знает, какие результаты остаются в рабочих файлах, какие подлежат хранению, а какие требуют дополнительного утверждения перед тем, как повлиять на регуляторную позицию.	повседневной деятельности
Оптимизированная	Управление ИИ рассматривается как непрерывная операционная дисциплина — сценарии использования пересматриваются по мере развития инструментов, обучение обновляется, а исключения отслеживаются и анализируются.	Регуляторные вопросы больше не реагируют на внедрение ИИ, они активно формируют его. Вопрос больше не звучит как «Разрешаем ли мы ИИ?», а как «При каких условиях использование ИИ улучшает регуляторную работу, не ослабляя подотчетность?»	Основная задача — поддержание дисциплины с течением времени

Измеряйте то, что важно

Модель зрелости оценивается по пяти областям — **Люди (People)**, **Процессы (Process)**, **Технологии (Technology)**, **Надзор (Oversight)** и **Управление доказательствами (Evidence Management)** — потому что управление редко дает сбой во всем сразу. Регуляторная команда может иметь доступ к безопасной, одобренной корпорации ИИ-платформе и при этом не иметь общего понимания того, что на самом деле означает «проверенный результат». Это пробел в области **Людей**, проявляющийся в среде с высокой технологической зрелостью, и единый общий балл зрелости никогда бы его не выявил.

Каждая область имеет свою специфику:

- **Люди** — Понимают ли пользователи, что инструменты могут, а что не могут делать? Знают ли проверяющие, как верифицировать результат ИИ? Назначен ли кто-то официально ответственным за управление ИИ в функции RA?
- **Процессы** — Документированы ли разрешенные и запрещенные сценарии использования? Является ли проверка стандартным этапом или личным выбором? Следуют ли исключения определенному пути?
- **Технологии** — Выбраны и одобрены ли инструменты для предполагаемого регуляторного использования? Контролируется ли доступ? Можно ли отследить результаты до конкретной сессии или версии?
- **Надзор** — Назначен ли конкретный проверяющий для результатов ИИ с более высоким уровнем риска? Ясны ли пути эскалации, когда качество результата вызывает сомнения?
- **Управление доказательствами** — Может ли организация объяснить, как использовался результат, созданный с помощью ИИ, как он проверялся и должен ли он сохраняться как часть рабочей или регулируемой документации?

Несколько прямых вопросов по каждой области обычно дают больше информации, чем любой формальный аудит. Если на эти вопросы трудно ответить, уровень зрелости становится очевидным.

Почему это важно именно сейчас

Регуляторные руководства движутся в одном направлении. Проект руководства FDA по использованию ИИ в принятии регуляторных решений подчеркивает важность целевого использования, риска результата и доказательств достоверности перед тем, как полагаться на них.¹ Акт ЕС об ИИ (**EU AI Act**) усиливает подотчетность на протяжении всего жизненного

цикла и человеческий надзор в сценариях с высокими последствиями.² Ничего из этого не является новой философией, так как это знакомая регуляторная логика, которая теперь прописывается эксплицитно в контексте ИИ. Структура управления рисками ИИ от NIST (**NIST AI Risk Management Framework**) добавляет пропорциональное управление, прослеживаемость и человеческий надзор в качестве основополагающих требований.⁴ Эти принципы: целевое использование, пропорциональность риска, человеческий надзор, прослеживаемость — уже глубоко знакомы специалистам по регуляторным вопросам. Модель зрелости переносит их во внутренний операционный контекст, где ИИ теперь используется повсеместно.

Ставки стали конкретными в апреле 2026 года, когда FDA выпустила предупредительное письмо (warning letter), содержащее раздел под заголовком «Ненадлежащее использование искусственного интеллекта в фармацевтическом производстве» — обозначение, которое ранее не встречалось в правоприменительной переписке агентства.³ Производитель использовал ИИ-агента для генерации СОП (стандартных операционных процедур), спецификаций и производственных записей без надлежащей проверки человеком. Когда следователи обнаружили, что валидация процесса не была проведена до начала распределения продукции, владелец ответил, что ИИ-инструмент не сигнализировал о необходимости этого требования. Предприятие было фармацевтическим, а не по производству медицинских изделий. Принцип идентичен: ИИ не несет регуляторной ответственности, её несут люди.

Переход от оценки к действию

Самой продуктивной отправной точкой является структурированная самооценка по пяти областям — по одному честному вопросу на каждую область, на который должен ответить кросс-функциональный коллектив. Документированы ли утвержденные сценарии использования? Назначен ли конкретный проверяющий для регуляторного контента, созданного с помощью ИИ? Может ли организация объяснить, что произошло, когда ИИ повлиял на регулируемое решение?

Модель зрелости полезна только в том случае, если она стимулирует действия. Организациям не нужно достигать самого высокого уровня, чтобы получить выгоду; им нужно лишь знать, где они находятся сейчас, и как выглядит следующий разумный шаг. Контрольный список готовности, сопоставленный с областями и содержащий критерии для каждого уровня

зрелости, служит практическим дополнением к этой структуре. Большинство организаций в процессе работы обнаруживают, что пробел находится не там, где они ожидали.

Заключение

Проблема, стоящая перед регуляторными отделами medtech-компаний, больше не заключается в том, будет ли ИИ использоваться внутри компании. Он уже используется. Более важный вопрос заключается в том, могут ли организации операционализировать это использование таким образом, чтобы оно было дисциплинированным, пропорциональным и соответствующим стандартам доказательности и подотчетности, которых требует регуляторная работа.

Полезным тестом, вдохновленным риск-ориентированным подходом, уже знакомым специалистам по регуляторным вопросам, является способность вашей организации ответить на три вопроса для любой регуляторной задачи с помощью ИИ:

- Каковы возможные последствия, если результат ИИ будет ошибочным?
- Какой уровень проверки человеком требуется для управления этим риском?
- Кто несет ответственность за конечный рабочий продукт?

Если ответ «нет», в вашей организации может быть внедрение ИИ, но еще нет управления ИИ. Разница между этими двумя результатами заключается не в инструменте. Она заключается в зрелости организации, которая его использует.

Ссылки

- S. Food and Drug Administration. Considerations for the use of artificial intelligence to support regulatory decision-making for drug and biological products. Draft Guidance. January 2025. Available at: <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/considerations-use-artificial-intelligence-support-regulatory-decision-making-drug-and-biological>
- European Parliament and Council of the European Union. Regulation (EU) 2024/1689 — Artificial Intelligence Act. Official Journal of the European Union. June 2024. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- S. Food and Drug Administration. Warning Letter to Purolea Cosmetics Lab (MARCS-CMS 722591). Issued April 2, 2026. Available at: <https://www.fda.gov/inspections-compliance-enforcement-and-criminal-investigations/warning-letters/purolea-cosmetics-lab-722591-04022026>
- National Institute of Standards and Technology. Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI 600-1).

July 2024. <https://www.nist.gov/itl/ai-risk-management-framework>

- Hussein, R., Zink, A., Ramadan, B., Howard, F. M., Hightower, M., Shah, S., et al. (2026). Advancing healthcare AI governance through a comprehensive maturity model based on systematic review. npj Digital Medicine, 9(1). <https://doi.org/10.1038/s41746-026-02418-7>. https://www.nature.com/articles/s41746-026-02418-7_reference.pdf

Перевод выполнен: 15.07.2026 | ai4med.ru

Машинный перевод. Рекомендуем сверять с оригиналом при клиническом использовании.