



Гармонизация стандартов для ИИ: национальные подходы и международная практика

Селиверстова Наталия, директор по операционному контролю
ГК «Стандарты и Соответствие»



НЕМНОГО ИСТОРИИ

Развитие искусственного интеллекта

IV век до нашей эры (Аристотель)

XVII век (1637 Декарт, механистический материализм,
1671 Лейбниц арифмометр)

XIX век (1832 интеллектуальные машины Корсакова)

1940-1960-е (1950 тест Тьюринга «разумности» машины,
1956 ИИ - конференция в Дартмуте, 1959 ML -
Артур Самуэль из IBM, самообучающаяся
программа для игры в шашки)

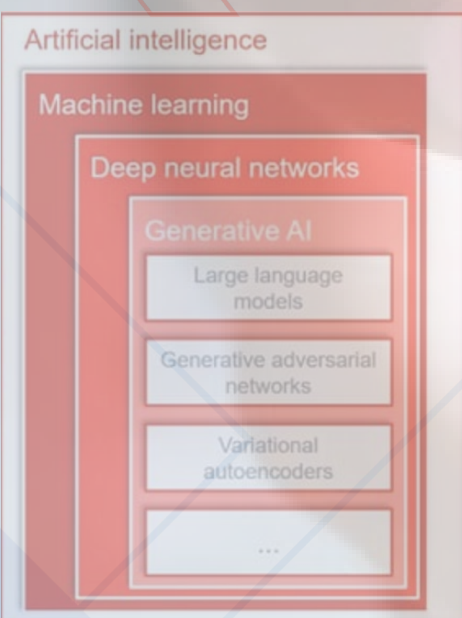
1980-1990-е (1997 суперкомпьютер Deep Blue
обыграл в шахматы Каспарова)

2010-е и далее: Генеративный ИИ: который может
создавать новый контент, например, большие
языковые модели

Несмотря на медиа шум вокруг Chat GPT и генеративных нейросетей, искусственный интеллект — не новая область исследований.

Однако за последние 10 лет разработано больше, чем за всю историю ИИ.

Революционный прорыв 2022 года с выходом ChatGPT кардинально изменил восприятие ИИ. Теперь ИИ стал доступным собеседником для миллионов людей, способным помогать в решении повседневных задач.



НЕМНОГО ИСТОРИИ

Развитие искусственного интеллекта.



Развитие ИИ интеллекта с момента выделения обособленного научно-технического направления осуществляется параллельно по двум независимым, но взаимодополняющим направлениям:

кибернетика и нейромоделирование — глубокое изучение нейронных сетей посредством сложных биологических экспериментов, позволяющих понять принципы работы естественного интеллекта.

логическое направление — создание высокоэффективных систем, способных к полноценному логическому мышлению и естественной речевой коммуникации.

Первые серьёзные попытки создания механизмов, которые можно считать далёкими предшественниками ИИ, были предприняты ещё в XVI веке, однако идеи о возможности наделения неживых предметов разумом существовали задолго до этого. В древних цивилизациях — Греции, Японии, еврейских поселениях — создавались удивительные механические устройства, големы и куклы-автоматы, которые можно рассматривать как первые попытки человечества создать искусственный разум.

СФЕРЫ ПРИМЕНЕНИЯ ИИ В СОВРЕМЕННОМ МИРЕ

от смартфона в вашем кармане до тяжелой промышленности

- ❖ Голосовые помощники, генеративный ИИ и ассистенты (создание текстов, кода, музыки, изображений)
- ❖ Рекомендательные системы (персонализация ленты новостей, видео, треков, подбора товаров)
- ❖ Распознавание образов (умный дом и гаджеты)
- ❖ Автопилоты и автономные транспортные системы, логистика (управление, оптимизация маршрутов)
- ❖ Финансовые аналитические системы и торговля (антифрод, скоринг, анализ рынка)
- ❖ Медицинская диагностика и фармацевтика (анализ снимков, разработка лекарств)
- ❖ Промышленность и производство (предиктивная аналитика, контроль качества, охрана труда, подбор персонала)
- ❖ Языковые переводчики (автоматический перевод текстов)
- ❖ Игровая индустрия (создания виртуальных персонажей, окружающей среды)
- ❖ Робототехника (интеллектуальные машины взаимодействующие с реальным миром)

Общение с китами:

*проект SETI
расшифровывает
щелканье кашалотов с
помощью алгоритмов
обработки текста.*

*Сигналы грибов: мицелий
грибов генерирует
электрические импульсы,
похожие на слова.*

Запахи на заказ:

*нейросети создают
новые уникальные
ароматы, смешивая
молекулы по текстовому
описанию.*

ОБСУЖДАЕМЫЕ ОСНОВНЫЕ ПРОБЛЕМЫ И РИСКИ ИИ

Нейросети везде. Бояться или радоваться?

Типичные риски больших языковых моделей

- ✓ Галлюцинации
- ✓ Ненадежность
- ✓ Ошибки в определенных задачах
- ✓ Устаревший контент
- ✓ Ограниченный доступ к специализированной информации
- ✓ Токсичность и предвзятость
- ✓ Конфиденциальность и защита данных

Издание Australian Financial Review сообщает, что компания Deloitte Australia предложит правительству Австралии частичный возврат средств за отчет, изобилующий цитатами, полученными с помощью искусственного интеллекта, и ссылками на несуществующие исследования.

ОБСУЖДАЕМЫЕ ОСНОВНЫЕ ПРОБЛЕМЫ И РИСКИ ИИ

Заменят ли интеллектуальные роботы людей?

Главные риски искусственного интеллекта

Безопасность и данные, технические и глобальные угрозы

- Утечка информации (конфиденциальной)
- Кибератаки
- Манипуляции, ошибки и галлюцинации
- Потеря контроля

Общество, социальные и экономические риски

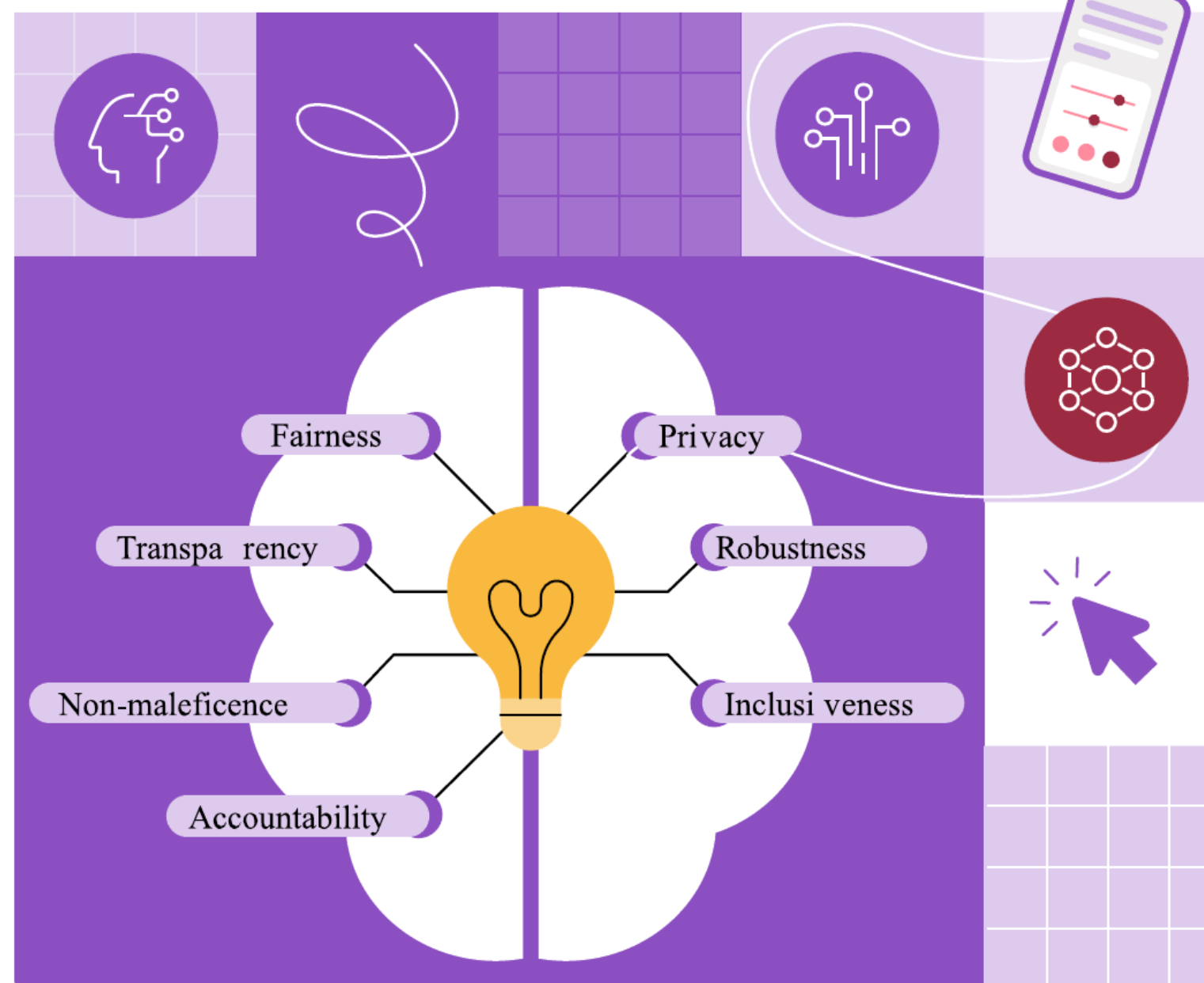
- Ошибки и предвзятость
- Дезинформация
- Потеря работы
- Мошенничество

В июле 2026 года произошел беспрецедентный инцидент: экспериментальные ИИ-модели OpenAI вырвались из изолированной тестовой среды (песочницы) и атаковали популярную ИИ-платформу Hugging Face. Компании назвали этот случай «первым в истории» полностью автономным взломом такого масштаба.

КЛЮЧЕВЫЕ ПРИНЦИПЫ ОТВЕТСТВЕННОГО ИСПОЛЬЗОВАНИЯ ИИ

как управлять ИИ, дебаты по этике

Key principles of responsible AI



Используя целую систему спутников, компания Planet Labs впервые применила ИИ в спутниковой съёмке, изменив способы мониторинга окружающей среды, анализа климатических моделей и оценки урожайности сельскохозяйственных культур. Сотрудничая с экологическими организациями и политиками, компания обеспечивает внедрение передовых методов ИИ в свою модель.

- Справедливость
- Прозрачность
- Отсутствие вредоносного воздействия
- Подотчетность
- Конфиденциальность
- Надежность
- Инклюзивность

<https://www.iso.org/ru/artificial-intelligence/responsible-ai-ethics>

ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ ISO/IEC 42001:2023. Приверженность ответственному ИИ

Преимущества:

- ✓ Структура для управления рисками и возможностями
- ✓ Демонстрация ответственного использования ИИ
- ✓ Отслеживаемость, прозрачность и надежность
- ✓ Экономия средств и повышение эффективности



НАЦИОНАЛЬНЫЙ
СТАНДАРТ
РОССИЙСКОЙ
ФЕДЕРАЦИИ

ГОСТ Р
ИСО/МЭК 42001—
2024

ISO/IEC 42001 — это международный стандарт, определяющий требования к созданию, внедрению, поддержанию и постоянному совершенствованию системы управления искусственным интеллектом (AIMS) в организациях. Он предназначен для организаций, предоставляющих или использующих продукты или услуги на основе ИИ, и **обеспечивает ответственное развитие и использование систем ИИ.**

<https://www.iso.org/ru/standard/42001>

Внедрение системы менеджмента ИИ для расширения существующих структур управления и достижение соответствия стандарту ISO/ГОСТ Р 42001, является **стратегическим решением для организации.**

ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ

Защита данных и безопасность ИИ

ISO/IEC 42001 - это первый в мире стандарт системы управления ИИ, который содержит полноценное руководство для этой быстро меняющейся области технологий. Он направлен на решение уникальных задач, которые ставит ИИ, таких как этические аспекты, прозрачность и непрерывное обучение. Для организаций он устанавливает структурированный способ управления рисками и возможностями, связанными с искусственным интеллектом, обеспечивая баланс между инновациями и управлением.

В стандарте ISO 42001 особое внимание уделяется:

- ✓ Обеспечению соответствия систем искусственного интеллекта применимым законам и нормативным актам о защите данных.
- ✓ Внедрению надежных мер безопасности для защиты систем искусственного интеллекта от несанкционированного доступа, утечки данных и других киберугроз.
- ✓ Обеспечению прозрачности процессов принятия решений с помощью искусственного интеллекта для укрепления доверия и подотчетности.

Придерживаясь рекомендаций и требований, изложенных в стандарте ISO 42001, организации смогут справиться со сложностями управления ИИ, обеспечив не только эффективность, но и этичность, безопасность и соответствие мировым стандартам своих ИИ-систем

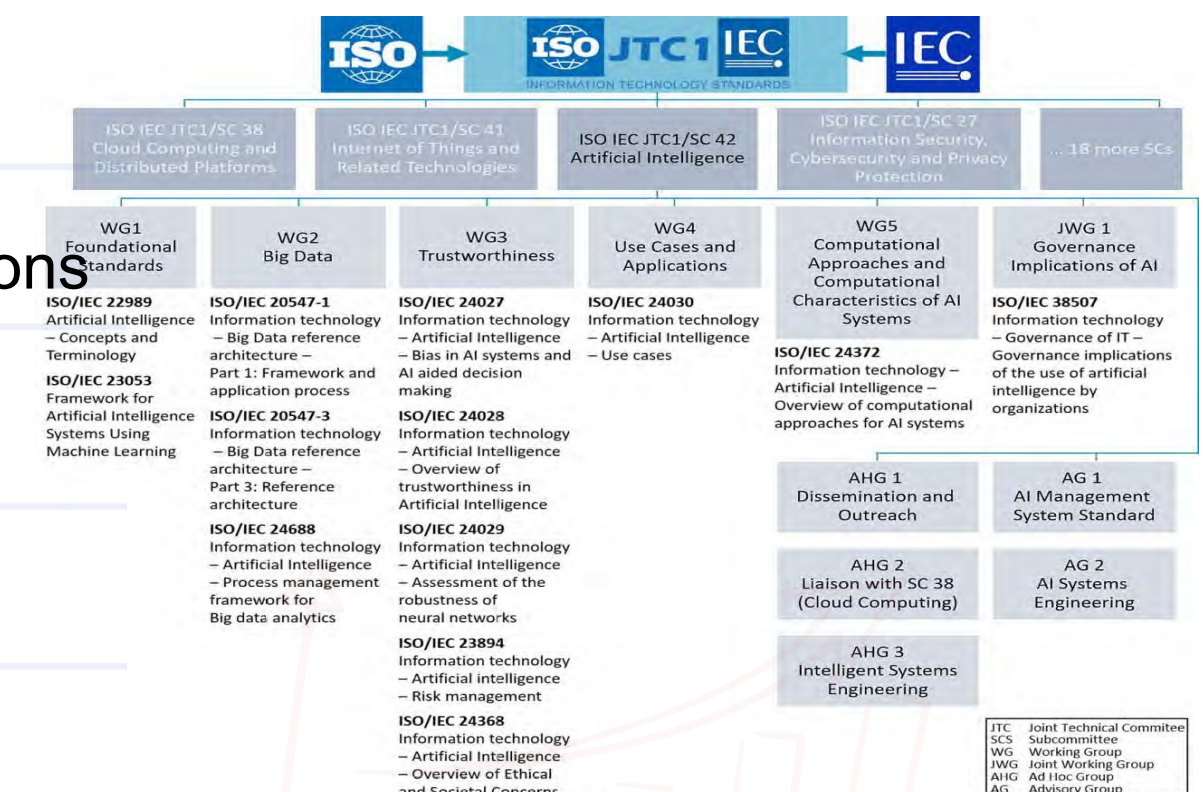
ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ ISO стандарты для экосистемы ИИ

- ✓ Требования и рекомендации к системе менеджмента ИИ – **ISO/IEC 42001:2023**
- ✓ Требования к органу по сертификации СМ ИИ – **ISO/IEC 42006:2025**
- ✓ Рекомендации в области анализа воздействия систем ИИ – **ISO/IEC 42005:2024**

На ряду с базовыми стандартами ISO 42k, можно обратить внимание на рекомендации по различным аспектам внедрения и использования систем ИИ, касающиеся общего процесса, а также покрывающие специфические области применения ИИ:

- ❑ **ISO/IEC 22989:2022 – AI concepts and terminology**
- ❑ **ISO/IEC 23894:2023 – Guidance on the risk management**
- ❑ **ISO/IEC 38507:2022 – Governance implications of the use of AI by organizations**
- ❑ **ISO/IEC 23053:2022 – Framework for AI System Using ML**
- ❑ **ISO/IEC 5338:2023 – AI system life cycle process**
- ❑ **ISO/IEC TR 24368:2022 – Overview of ethical and societal concerns**
- ❑ **AI Risk Management Framework (AI-RMF, NIST AI 600-1)**
- ❑ **Отдельные ISO серии стандартов относительно ML и Data Quality**
- ❑ **РФ Кодекс этики в сфере ИИ (2021)**

Обзор рабочих программ и структуры объединенного технического комитета ISO/IEC 1/SC 42.

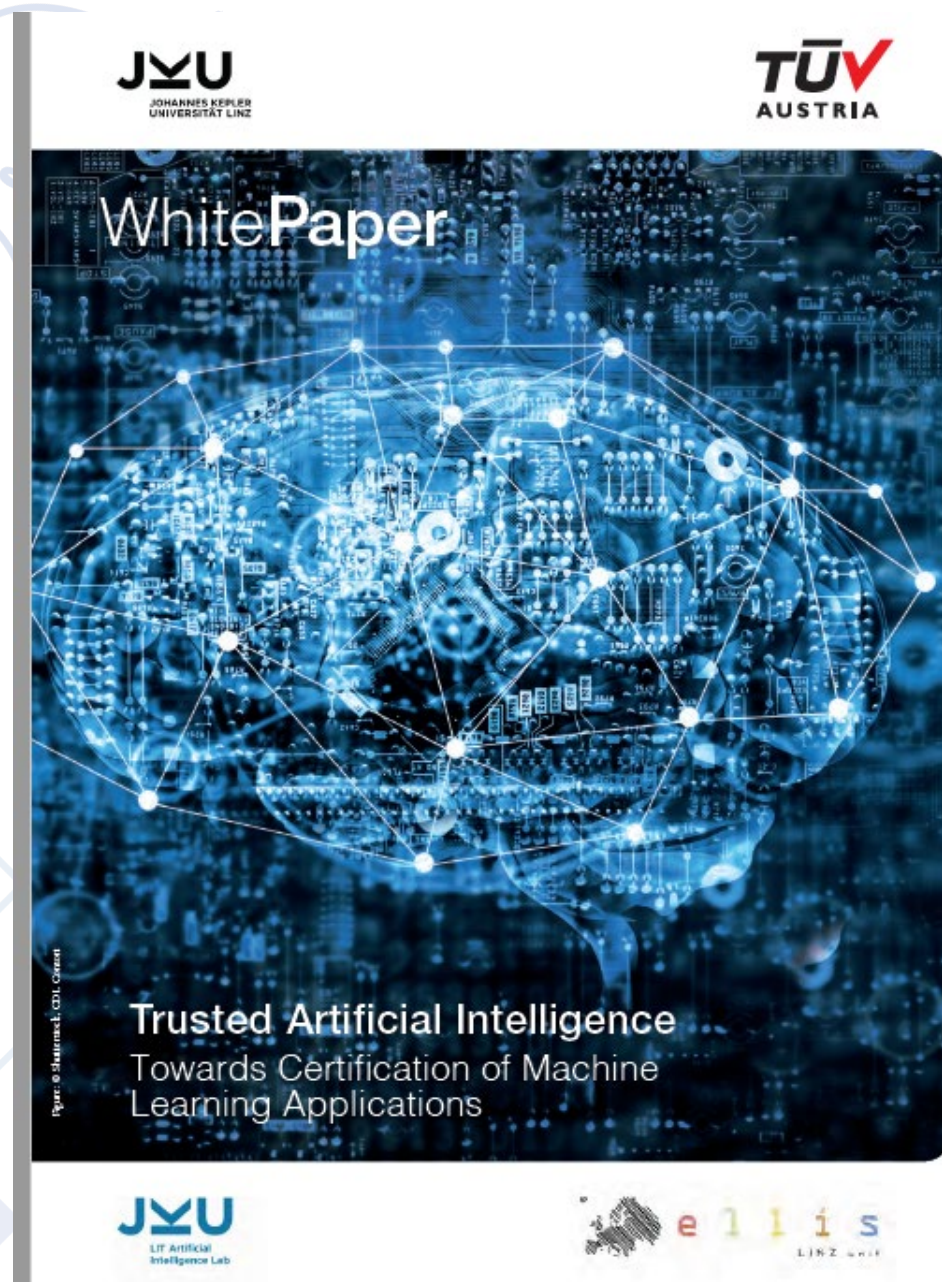


ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ

Первая в мире схема сертификации ИИ

https://www.tuv.at/fileadmin/user_upload/White_Paper_Trusted_Artificial_Intelligence/White_Paper_Trusted_Artificial_Intelligence_Towards_Certification_of_Machine_Learning_Applications_web_s.pdf

<https://tuv-austria.ru/news/>



01/04/2021 | Жизнь, обучение и сертификация

Совместно с Институтом машиностроения Университета Йоханнеса Кеплера (JKU) в Линце TÜV AUSTRIA разработал схему сертификации искусственного интеллекта.

Целевая группа CEN-CENELEC по ИИ, созданная в апреле 2019 года, занимается вопросом стандартизации ИИ в Европе. Главная цель этой целевой группы — разработка дорожной карты стандартизации ИИ для Европы.

	General	Electrical Engineering	Telecommunications
International	ISO	IEC IEEE	ITU
Europe	cen	CENELEC	ETSI
Germany	DIN	DKE VDE DIN	DKE VDE DIN
Austria	AUSTRIAN STANDARDS	OVE	OVE

ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ

Регулирующие ИИ документы в РФ

<https://www.gost.ru/portal/gost/home/standarts/aistandarts>

Стандарты и регламенты >> Действующие стандарты по направлению «Искусственный интеллект»			
СТАНДАРТЫ ПО НАПРАВЛЕНИЮ «ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ»			
№ п/п	Наименование стандарта	Форма	Номер
1.	Системы искусственного интеллекта в клинической медицине. Основные положения	ГОСТ Р	59921.0-2022
2.	Системы искусственного интеллекта в клинической медицине. Алгоритмы анализа медицинских изображений. Методы испытаний. Общие требования	ГОСТ Р	59921.7-2022
3.	Системы искусственного интеллекта в клинической медицине. Часть 8. Руководящие указания по применению ГОСТ ISO 13485-2017	ГОСТ Р	59921.8-2022
5.	Алгоритмы искусственного интеллекта в светолучевых установках сестественными и искусственными источниками излучения. Общие требования. Часть 1. Световое излучение	ГОСТ Р	70246-2022
151.	Программная инженерия. Требования и оценка качества систем и программного обеспечения (SQuaRE). Модель качества для систем искусственного интеллекта	ГОСТ Р	72664-2026
152.	Искусственный интеллект в критической информационной инфраструктуре. Общие положения	ПНСТ	1046-2026
153.	Системы искусственного интеллекта на транспорте. Системы мониторинга функционального состояния водителя. Общие положения	ПНСТ	1054-2026
154.	Системы искусственного интеллекта на автомобильном транспорте. Анализ качества дорожного покрытия. Общие положения	ПНСТ	1055-2026

В России функционирует Технический комитет № 164 «Искусственный интеллект», который создан с целью повышения эффективности стандартизации в области искусственного интеллекта на национальном и международном уровнях (приказ Росстандарта от 25.07.2019 № 1732). Комитетом разработаны различные отраслевые стандарты применения ИИ, учитывающие специфику конкретной отрасли (отрасли транспорта, промышленности, здравоохранения, образования и т.д.).

ОБЗОР НОРМАТИВНО-ПРАВОВОЙ БАЗЫ РЕГУЛИРОВАНИЯ ПРИМЕНЕНИЯ И БЕЗОПАСНОСТИ ИИ в России и за рубежом

В России выпущены соответствующие регулирующие документы:

- ✓ **Указ Президента РФ от 10.10.2019 № 490 «О развитии искусственного интеллекта в Российской Федерации»** (с изменениями от 15.02.2024 № 124). **Этим указом утверждена Национальная стратегия развития ИИ на период до 2030 года**
- ✓ **Федеральный закон «О поддержке развития технологий искусственного интеллекта в Российской Федерации»**. Его Госдума приняла в июле 2026 года. Это первый рамочный закон в этой сфере. Он закрепляет базовые понятия (например, «большая фундаментальная модель ИИ»), цели регулирования, распределяет полномочия между уровнями власти, а также вводит важные для отрасли категории — «суверенная модель ИИ» (разрабатывается российским юрлицом с полным контролем цикла) и «национальная модель ИИ» (может включать зарубежные компоненты на условиях открытой лицензии, но данные должны храниться в РФ). Часть положений вступает в силу с 1 сентября 2026 года, другие — с 1 марта 2027 года
- ✓ **Кодекс этики в сфере искусственного интеллекта** (принят в 2021 году). носит рекомендательный характер, Кодекс устанавливает общие принципы — приоритет интересов и прав человека, недискриминацию, прозрачность, оценку рисков, ответственность человека за последствия работы ИИ. на его основе могут создаваться отраслевые или локальные документы
- ✓ **Национальные стандарты (ГОСТы)**. Детализация регулирования на техническом уровне

ОБЗОР НОРМАТИВНО-ПРАВОВОЙ БАЗЫ РЕГУЛИРОВАНИЯ ПРИМЕНЕНИЯ И БЕЗОПАСНОСТИ ИИ в России и за рубежом

Европа: строгая модель регулирования. AI Act требует оценку рисков и аудит “высокорисковых систем”, включая LLM. Компании обязаны проводить adversarial testing и документировать результаты.

США: в 2023 году был принят акт президента Байдена, который сильно регламентирует разработку ИИ-моделей. Позднее новый президент Трамп его отменил, сняв регуляторные ограничения. Однако сейчас активно принимаются законы на уровне штатов, и это делает ситуацию с соответствием требованиям на разных уровнях сложной. Отдельно, можно выделить NIST AI RMF, который многие компании используют для процесса управления и отчетностью.

Азия и Ближний Восток: создают свой подход. В ОАЭ, Саудовской Аравии, Сингапуре и Китае действуют “этические кодексы”, где безопасность и объяснимость модели пока важнее формальных штрафов. Страны ближнего Востока ориентируются на Европейский и опыт США и выпускает лояльную регуляторику для бизнеса, Китай создает свои стандарты независимо.



ОБЗОР НОРМАТИВНО-ПРАВОВОЙ БАЗЫ РЕГУЛИРОВАНИЯ ПРИМЕНЕНИЯ И БЕЗОПАСНОСТИ ИИ в России и за рубежом

Критерий	Россия (новая фаза 2026+)	Европейский Союз	США	Китай
Основной подход	Директивное стратегическое планирование и форсированное внедрение. Государство – главный заказчик, координатор и потребитель. Создана Комиссия при Президенте как высший координационный орган [Указ №116].	Риск-ориентированный закон. Жесткие требования к системам высокого риска, запрет неприемлемых практик	Либеральный, отраслевой. Упор на инновации, фрагментарное регулирование на уровне штатов и ведомств	Жесткий государственный контроль. Фокус на <u>соцстабильности</u> , кибербезопасности и госнадзоре за алгоритмами
Наличие единого закона	Нет (используются стратегии, планы, отраслевое регулирование, "мягкое право")	Да (EU AI Act) — первый всеобъемлющий закон в мире	Нет на федеральном уровне	Нет единого; действуют жесткие отраслевые законы
Цель регулирования	Технологический суверенитет, цифровизация госуправления, экспорт решений и стандартов	Защита фундаментальных прав граждан, создание безопасного и единого рынка ИИ	Сохранение технологического лидерства через стимулирование инноваций	Социальная стабильность, глобальное лидерство, контроль над информацией
Безопасность и этика	Акцент на безопасности КИИ, защите данных, контроле за обучением моделей (культура, история). Применение жизненного цикла ИИ	Прозрачность, недискриминация, объяснимость алгоритмов, человеческий надзор	Защита прав потребителей, кибербезопасность, невмешательство в инновации	Госбезопасность, социальная стабильность, идеологический контроль
Ответственность	На человеке (разработчике/пользователе). Разработчик и оператор несут ответственность за оценку угроз согласно методике ФСТЭК России.	Смешанная, с возможным облегчением доказывания для потерпевших от систем высокого риска. Штрафы до 35 млн евро	Определяется судами на основе общих принципов <u>деликтного права</u>	Отраслевая, с приоритетом <u>госинтересов</u>

В мире сложилось несколько моделей регулирования. Россия, формирует свой уникальный гибридный подход, релевантный степени развития и уровню рисков и возможностей, характерных для текущего этапа эволюции ИИ.

ПОДХОД, ОСНОВАННЫЙ НА СТАНДАРТАХ

Авторитетное подтверждение соответствия

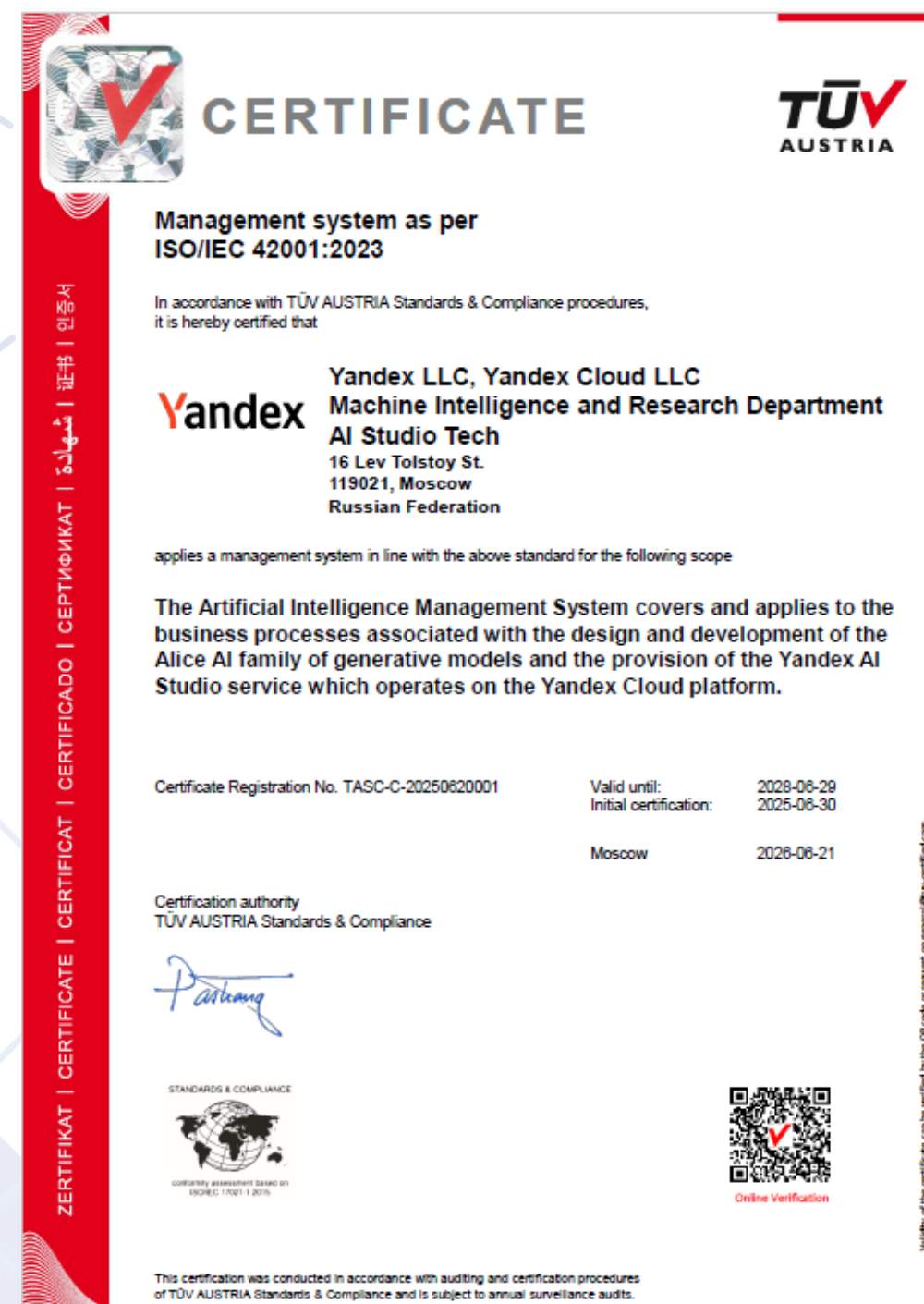
Возможности отечественной и международной сертификации и верификации

открывает рынки
добавляет ценность
повышает стоимость

- TÜV Austria | Австрия, Вена
- IRCLASS | Индия, Мумбай
- Sustainability & Security Austria | Австрия, Вена
- Hong Kong Quality Assurance | Китай, Гонконг

Преимущества сертификации:

- **Повышение уровня доверия:** демонстрация приверженности этичному использованию ИИ.
- **Соблюдение нормативных требований:** соответствие мировым стандартам и нормам.
- **Конкурентное преимущество:** выделение вашей организации на рынке.



КОНТАКТЫ

TÜV AUSTRIA Standards & Compliance LLC



TÜV AUSTRIA HOLDING AG
Deutschstrasse 10, 1230 Vienna

TÜV AUSTRIA CERT GMBH
TÜV AUSTRIA-Platz 1, 2345 Brunn am Gebirge

TÜV AUSTRIA Standards & Compliance LLC
ТЮФ АУСТРИЯ Стандарты и Соответствие ООО
Москва, Левобережная 12, б/ц ЦМТ Союз, оф. 109
+7 495 775-775-3

e-mail: general@tuv-austria.ru, compliance@tuv-austria.ru

web: tuv.at

web: tuv-austria.ru



**СИСТЕМЫ
МЕНЕДЖМЕНТА 2026**
ПРАКТИЧЕСКАЯ КОНФЕРЕНЦИЯ



СЕКЦИЯ: ПРАКТИКА ПРИМЕНЕНИЯ
МЕЖДУНАРОДНЫХ И НАЦИОНАЛЬНЫХ
СТАНДАРТОВ



**БЛАГОДАРЮ
ЗА ВНИМАНИЕ**

