

МОШЕННИЧЕСТВО ПРИ ЦИФРОВОЙ ИДЕНТИФИКАЦИИ:

**угрозы, масштабы, технологии противодействия и научный
задел**

Состояние на 2026 год

Версия 2.0 - редакция для внешнего распространения

*(закреты фактологический, методологический и нормативный уровни качества; нормативный раздел
сверен с первичными актами)*

Август 2026 года

Содержание

Резюме для руководства	4
1. Методология.....	5
1.1. Предмет и границы исследования.....	5
1.2. Классификация источников и уверенность в выводах.....	5
1.3. Паспорт показателя и несопоставимость метрик.....	6
1.4. Что установлено, что вероятно, что не доказано	7
2. Часть I. Масштабы и динамика угроз (2024–2026)	8
2.1. Карта атак по жизненному циклу цифровой идентичности	8
2.2. Глобальные ориентиры (уровень A1, оценочные значения)	8
2.3. Соединённые Штаты Америки (уровень A1).....	9
2.4. Великобритания (уровень A1).....	9
2.5. Европейский союз (уровень A1).....	9
2.6. Российская Федерация.....	10
2.7. Данные сетей верификации (уровни B/D).....	10
2.8. Экономика атакующего.....	11
2.9. Человеческий фактор.....	11
2.10. Показательный инцидент: дело Agur (уровень A4, рамка отдельного случая)	12
3. Часть II. Технологии и практики противодействия	13
3.1. Пятислойная модель защиты	13
3.2. Независимые государственные испытания DHS S&T RIVTD (уровень A3, по первичным документам).....	13
3.3. Прочие независимые испытания (уровень A3)	15
3.4. Рынок решений: доказательная матрица вместо перечня.....	16
3.5. Регуляторный контекст (уровень A2, сверено с первичными актами).....	17
3.6. Тренды на горизонте 12–24 месяцев.....	18
4. Часть III. Научный задел.....	20
4.1. Организация по исследовательским вопросам.....	20
4.2. Ключевые публикации и эталонные наборы.....	20
4.3. Сводный вывод по научному разделу.....	21
5. Рекомендации для организаций	22
5.1. Триггеры приоритизации от собственных данных.....	22
5.2. Этапный план	22
5.3. Обязательные вопросы поставщику до пилота или закупки	23
6. Ограничения исследования	24
7. Реестр источников	25
Уровень A1. Официальная статистика и административные данные.....	25
Уровень A2. Нормативные акты и официальные руководства	25
Уровень A3. Государственные технические испытания (первичные публикации)	26
Уровень A4. Расследованные инциденты.....	26
Уровень B. Научные публикации (препринты помечены)	26
Уровень C. Отраслевые исследования.....	27
Уровень D. Отчёты поставщиков	27
Приложение А. Основные термины и метрики	28
Приложение Б. Официальные таблицы результатов RIVTD (MdTF).....	29
Б.1. Трек 2: сопоставление селфи с документом (16 систем; 1 633 добровольца)	29
Б.2. Трек 3: активные подсистемы детекции презентационных атак (6 систем; 660 участников)	29

Б.3. Трек 3: пассивные подсистемы детекции презентационных атак (15 систем)	30
Приложение В. Юридическая верификация нормативного раздела	31
В.1. Верификационная таблица юридически значимых утверждений	31
В.2. Матрица применимости регулирования по типам систем	33
В.3. Оговорка.....	33

Резюме для руководства

Настоящее исследование посвящено мошенничеству при цифровой идентификации в период 2024–2026 годов. В отличие от отраслевых отчётов поставщиков решений, исследование построено на триангуляции независимых наблюдательных контуров: официальной статистики правоохранительных органов и регуляторов, государственных технических испытаний, данных транзакционных и верификационных сетей, научных публикаций и отдельных расследованных инцидентов. Каждому источнику присвоен уровень доказательности, каждому ключевому показателю - «паспорт», фиксирующий числитель, знаменатель, период, географию и ограничения.

Первое. Масштаб угрозы независимо подтверждается несколькими не связанными между собой контурами наблюдения, однако единая «глобальная цифра потерь» остаётся оценкой, а не измерением. ИНТЕРПОЛ в оценке угроз финансового мошенничества (второе издание, март 2026 г.) оценивает глобальные потери за 2025 год более чем в 442 млрд долларов США, присваивая глобальному риску уровень «высокий» и ожидая существенной эскалации в ближайшие три–пять лет вследствие доступности технологий искусственного интеллекта. Эта оценка сопоставима по порядку величины с фигурировавшей в отчёте Liminal + Unico цифрой «свыше 400 млрд долларов», но, в отличие от последней, имеет институционального автора и раскрытый методологический контекст.

Второе. Атаки дешевеют и усложняются одновременно. По оценке ИНТЕРПОЛА, мошенничество с применением искусственного интеллекта в 4,5 раза прибыльнее традиционных методов; на криминальных площадках наблюдаются предложения «дипфейк-как-услуга» и пакетов синтетических личностей по минимальным ценам от единиц долларов. Сети верификации фиксируют сходное направление: рост инъекционных атак, дипфейков и синтетических личностей при снижении доли примитивных презентационных атак. Направление тенденции подтверждено; абсолютные значения между сетями несопоставимы.

Третье. Идентичность стала главным каналом вреда. В Великобритании 72% всех записей Национальной базы мошенничества Cifas за 2025 год связаны с мошенничеством с идентичностью и захватом продуктов и счетов; при этом само мошенничество с идентичностью снизилось на 3% - преступники смещаются к захвату действующих учётных записей, включая рост несанкционированных подмен SIM-карт на 38%. Это подтверждает центральный методологический выбор исследования: рассматривать не только первичную проверку личности, но полный жизненный цикл цифровой идентичности - от регистрации до восстановления доступа и межорганизационного использования.

Четвёртое. Ручная проверка не является надёжным контролем против современных подделок, а однослойная техническая проверка недостаточна: первичные данные государственных испытаний DHS S&T (Maryland Test Facility) показывают, что в худшем сценарии отдельные пассивные подсистемы детекции презентационных атак пропускали до 100% атак отдельных видов, тогда как лучшие сочетали ошибку на добросовестных предъявлениях не выше 0,5% с ошибкой на атаках не выше 3,3%. Работающий ответ - многослойный контур: доверенный захват и аттестация устройства, детекция инъекций, анализ генеративных артефактов, поведенческие и сетевые сигналы, непрерывная переоценка риска, а на инфраструктурном уровне - криптографически проверяемые

удостоверения. Пороги эффективности должны определяться независимыми испытаниями на репрезентативной для организации выборке, а не маркетинговыми показателями поставщиков.

Пятое. Регуляторное давление синхронизировано с технологическим: обязательства Регламента ЕС об искусственном интеллекте для высокорисковых систем (включая биометрическую идентификацию по приложению III) применяются со 2 декабря 2027 года; четвертая редакция руководства NIST SP 800-63 (июль 2025 г.) закрепляет риск-ориентированную модель, фишинго-устойчивую аутентификацию и цифровые кошельки; в Российской Федерации ужесточаются требования к работе с биометрией через Единую биометрическую систему.

1. Методология

1.1. Предмет и границы исследования

Под мошенничеством при цифровой идентификации в настоящем исследовании понимаются противоправные действия, в которых механизмом атаки выступает создание, подделка, присвоение, компрометация или злоупотребление цифровой идентичностью. Объект исследования - полный жизненный цикл цифровой идентичности: первичное подтверждение личности (identity proofing), выдача учётных данных, аутентификация, восстановление доступа, подтверждение критических операций, непрерывная оценка риска, блокировка и отзыв, межорганизационный обмен сигналами и оспаривание ошибочных решений.

Платёжное, инвестиционное и иное финансовое мошенничество включается в рассмотрение только в той части, в которой цифровая идентичность, выдача себя за другое лицо, захват учётной записи, синтетическая личность или компрометация учётных данных являются механизмом атаки. Данные о совокупных финансовых потерях используются как макроориентиры и не интерпретируются как статистика мошенничества с идентификацией.

1.2. Классификация источников и уверенность в выводах

Каждому источнику присвоен уровень доказательности. Уровень А разделён на четыре подуровня, поскольку официальная статистика, нормативные акты, лабораторные испытания и расследованные инциденты представляют собой разные виды доказательств: нормативный акт не подтверждает эффективность технологии, а лабораторный тест не измеряет финансовые потери.

Уровень	Содержание	Примеры
A1	Официальная статистика и административные данные	FBI IC3; Банк России; UK Finance; Cifas; ЕБА/ЕЦБ; ИНТЕРПОЛ
A2	Законы, нормативные акты и официальные разъяснения	Регламент (ЕС) 2024/1689; eIDAS 2.0; 572-ФЗ; NIST SP 800-63-4
A3	Государственные технические испытания и эталонные тесты	DHS S&T RIVTD (первичные публикации MdTF); NIST FATE-PAD (NISTIR 8491); NIST FATE MORPH
A4	Судебные документы и официально расследованные инциденты	Дело Arup (полиция Гонконга)

Уровень	Содержание	Примеры
B	Рецензируемые научные публикации; препринты маркируются отдельно	IEEE TPAMI; Information Fusion; материалы CVPR/ICCV; arXiv (с пометкой)
C	Отраслевые исследования и опросы с раскрытой методикой	TransUnion; LexisNexis Risk Solutions; Deloitte
D	Отчёты и данные поставщиков решений (заинтересованная сторона)	Entrust; Sumsb; iProov; Smile ID; Group-IB; Shufti; deepidv; самоидентификация участников анонимизированных испытаний
E	СМИ и экспертные публикации	Используются для обнаружения случаев и поиска первичных документов, не как основание количественного вывода

Помимо надёжности источника, для каждого ключевого вывода фиксируется уверенность: высокая, средняя, низкая или прогнозная. Надёжный источник может содержать метрику, ограниченно применимую к исследуемому вопросу.

1.3. Паспорт показателя и несопоставимость метрик

Большинство фигурирующих в отрасли цифр несопоставимы напрямую. Одни источники считают попытки, другие - успешные атаки, третьи - потенциальную экспозицию, четвёртые - фактические возмещённые убытки. Доля (в процентах) и абсолютная экспозиция (в денежном выражении) - разные измерения; проценты роста критически зависят от базы отсчёта и определения. Для ключевых показателей исследования составлены «паспорта», примеры которых приведены ниже.

Показатель	Числитель	Знаменатель	Период, география	Ограничение
«Дипфейк - каждая пятая биометрическая атака» (Entrust)	Попытки, классифицированные сетью Entrust/Onfido как дипфейк	Обнаруженные биометрические атаки в той же сети (не все верификации и не всё мошенничество)	Сент. 2024 - сент. 2025; >1 млрд верификаций, 195 стран	Клиентская база одного поставщика; классификация выполнена самим поставщиком
\$893,3 млн потерь от ИИ-мошенничества (IC3)	Скорректированные потери по 22 364 обращениям, где заявители указали применение ИИ	Все обращения в IC3 за 2025 г. (1 008 597)	2025 г.; США	Только распознанные заявителями случаи; не эквивалент статистики дипфейков
29,3 млрд руб. (Банк России)	Объём операций без добровольного согласия клиентов	Операции, отражённые в отчётности кредитных организаций	2025 г.; Российская Федерация	Макроориентир хищений при переводах; вклад механизмов цифровой идентификации не выделяется

Показатель	Числитель	Знаменатель	Период, география	Ограничение
72% записей связаны с идентичностью (Cifas)	Записи категорий «identity fraud» и «facility takeover»	Все записи Национальной базы мошенничества (444 993)	2025 г.; Великобритания	База участников одной системы обмена; не национальная генеральная совокупность
\$442 млрд глобальных потерь (ИНТЕРПОЛ)	Оценка совокупных потерь от финансового мошенничества	Мировая экономика	2025 г.; глобально	Экспертная оценка, а не аудированная сумма; охватывает всё финансовое мошенничество, не только идентификационное

1.4. Что установлено, что вероятно, что не доказано

Статус	Утверждения
Установлено (сходимость нескольких независимых контуров уровней A1/A3 и B/D)	Рост числа и доли инъекционных атак и дипфейков в потоках верификации; рост синтетических личностей; смещение преступной активности от первичной регистрации к захвату учётных записей и восстановлению доступа; неспособность невооружённой ручной проверки служить надёжным контролем; существенный и документированный первичными данными разброс качества между поставщиками по итогам государственных испытаний
Вероятно (согласованные данные уровней C/D при отсутствии подтверждения уровня A)	Кратное удешевление организации сложной атаки; опережающий рост «агентного» автоматизированного трафика; лидерство Латинской Америки по доле синтетических личностей как опережающий индикатор для других рынков
Не доказано / прогнозно	Точная величина глобальных потерь (все имеющиеся значения - оценки); прогнозы кратного роста дипфейк-мошенничества в 2026 г. (Shufti) и удвоения потерь финансовых организаций к 2030 г.; тезис о скором глобальном доминировании синтетических идентичностей

2. Часть I. Масштабы и динамика угроз (2024–2026)

2.1. Карта атак по жизненному циклу цифровой идентичности

Безопасная первичная регистрация не предотвращает последующий захват учётной записи, мошенническое восстановление доступа, подмену SIM-карты или устройства, перехват сессии, принуждение настоящего пользователя и использование легитимной учётной записи в качестве «дропа». Модель атак по этапам жизненного цикла приведена ниже.

Этап	Основные атаки
Регистрация	Поддельный документ; синтетическая личность; дипфейк; морфинг лица
Проверка документа	Замена полей; полностью сгенерированный документ; повторное предъявление (replay); инъекция экрана
Биометрическая проверка	Фотография; маска; видеоповтор; дипфейк; виртуальная камера; инъекция потока
Выдача учётных данных	Перехват канала; подмена устройства; компрометация поставщика учётных данных
Вход	Фишинг; перебор утёкших учётных данных (credential stuffing); кража сессии; «усталость» от многофакторных запросов
Восстановление доступа	Подмена SIM-карты; выдача себя за клиента перед службой поддержки; подмена адреса электронной почты
Операция	Санкционированные жертвой платежи под воздействием (APP); дипфейк-распоряжение; мошенничество с удалённым доступом
Длительное использование	Выращивание учётных записей; «спящие» аккаунты; сети подставных счетов (мулов)
Межорганизационное использование	Одна личность, одно устройство или один актор в десятках организаций

2.2. Глобальные ориентиры (уровень A1, оценочные значения)

ИНТЕРПОЛ, «Global Financial Fraud Threat Assessment», второе издание, март 2026 г.: глобальные потери от финансового мошенничества за 2025 год оцениваются более чем в **442 млрд долларов США**; общий глобальный риск классифицирован как «высокий», ожидается существенная эскалация в ближайшие три–пять лет вследствие доступности технологий искусственного интеллекта и низких барьеров входа. Мошенничество с применением ИИ оценивается как **в 4,5 раза более прибыльное**, чем традиционные методы; «агентные» системы ИИ способны автономно планировать и исполнять полные кампании - от разведки до требования выкупа. На криминальных рынках зафиксированы предложения «дипфейк-как-услуга», инструментов клонирования голоса и пакетов синтетических личностей. С 2024 года число связанных с мошенничеством уведомлений ИНТЕРПОЛА выросло на 54%; организация поддержала более 1 500 транснациональных дел с возвращёнными активами на 1,1 млрд долларов. Значение 442 млрд долларов следует трактовать как институциональную оценку порядка величины, охватывающую всё финансовое мошенничество, а не только идентификационное.

Данная оценка выполняет в настоящем исследовании ту же функцию, что и цифра «свыше 400 млрд долларов» в отчёте Liminal + Unico, однако, в отличие от последней, имеет

раскрытого институционального автора, дату и методологический контекст, что переводит её из категории неподтверждаемых агрегатов в категорию проверяемых ориентиров.

2.3. Соединённые Штаты Америки (уровень A1)

FBI IC3, годовой отчёт за 2025 год: впервые за 25-летнюю историю центра выделена категория мошенничества с применением искусственного интеллекта. В 2025 году IC3 получил **22 364 обращения, содержащих сведения об использовании ИИ**, с совокупными скорректированными потерями более **893,3 млн долларов**. Показатель основан на обращениях заявителей, не охватывает нераспознанные случаи и не должен интерпретироваться как полная оценка мошенничества с применением ИИ; это один из наиболее значимых государственных количественных ориентиров, но не единственный государственный источник в мире. Крупнейшая подкатегория - инвестиционное мошенничество с ИИ-компонентом (около 632 млн долларов). Общие зарегистрированные потери от киберпреступности достигли 20,877 млрд долларов (+26% к 2024 г.) при 1 008 597 обращениях.

TransUnion (уровень C): экспозиция кредиторов США по синтетическим личностям оценена в **3,3 млрд долларов** по итогам 2024 года; в автокредитовании - около 2,0 млрд долларов только за первое полугодие 2024 года, что вдвое превышает сегмент банковских карт. Показатель отражает потенциальную экспозицию по открытым счетам, а не фактические потери.

2.4. Великобритания (уровень A1)

UK Finance, «Annual Fraud Report 2026» (июнь 2026 г.): в 2025 году через платёжное мошенничество похищено **1,28 млрд фунтов** (+4% к предыдущему году, второй год роста подряд); предотвращено несанкционированных операций ещё на 1,68 млрд фунтов. Структура динамики неоднородна: несанкционированные потери (преимущественно карточное мошенничество) снизились на 5% до 703,4 млн фунтов при росте числа случаев до 3,81 млн (+11%), тогда как санкционированные жертвой платежи под воздействием злоумышленника (APP) выросли на 19% до 576,4 млн фунтов. Мошенничество с выдачей себя за банк или полицию снизилось: потери -12%, число случаев -11%. 66% случаев APP начинались на интернет-платформах (32% потерь), 17% - через телекоммуникационные каналы (28% потерь). Возмещено жертвам APP 354,3 млн фунтов (61% потерь).

Cifas, «Fraudscape 2026»: в Национальную базу мошенничества за 2025 год внесено **444 993 записи** - максимум за всю историю наблюдений (+6%); участники фиксировали более 1 200 случаев ежедневно и предотвратили оценочно 2,4 млрд фунтов потерь. **72% всех записей связаны с мошенничеством с идентичностью и захватом продуктов и счетов (facility takeover)**. При этом само мошенничество с идентичностью снизилось на 3% (242 003 записи; 54% базы), что отражает смену тактики, а не снижение вреда: число захватов выросло на 6% (свыше 78 000 записей; около 90% приходится на мобильную связь, интернет-торговлю и кредитные карты), несанкционированные подмены SIM-карт выросли на **38%**, злоупотребление собственными счетами - на 43% (106 497 записей), зафиксировано более 22 000 записей по новой категории подставных счетов. Данные относятся к участникам одной системы обмена и не являются национальной генеральной совокупностью.

Совокупно британские данные дают важный контрпример упрощённому тезису «все виды мошенничества постоянно растут»: отдельные категории снижаются, при этом преступная активность перетекает в наименее защищённые этапы жизненного цикла - прежде всего в захват действующих учётных записей и восстановление доступа.

2.5. Европейский союз (уровень A1)

По совместной оценке Европейского банковского управления и Европейского центрального банка, совокупная стоимость выявленного платёжного мошенничества в Европейской экономической зоне увеличилась с 3,4 млрд евро в 2022 году до 4,2 млрд евро в 2024 году. Регуляторы считают усиленную аутентификацию клиента (SCA) эффективной, но фиксируют адаптацию преступников - смещение к социальной инженерии и операциям, которые пользователь подтверждает сам. Это согласуется с британскими данными о росте APP-мошенничества.

2.6. Российская Федерация

Банк России, «Обзор операций, совершённых без добровольного согласия клиентов» (уровень A1): объём таких операций составил **27,5 млрд рублей в 2024 году** (+74,4% к 2023 г.) и **29,3 млрд рублей в 2025 году** (+6,4% при росте числа операций на 31,2%). Доля возврата средств клиентам снизилась с 9,9% (2 713,6 млн руб.) в 2024 году до 5,9% (1 731,3 млн руб.) в 2025 году. Объём предотвращённых банками операций в 2025 году - около 13,9 трлн рублей. Показатели Банка России характеризуют общий контур хищений при переводах денежных средств; они используются в настоящем исследовании как макроориентир и не позволяют выделить вклад дипфейков, синтетических личностей или иных конкретных механизмов цифровой идентификации. Выделенной официальной статистики по дипфейкам Банк России не публикует.

Отраслевые оценки (уровни C/E, несопоставимые методики): «Информзащита» - около 5 700 дипфейк-атак в финансовом секторе за 2024 год (+13%); Сбербанк - рост числа дипфейков в 26 раз с начала 2025 года при среднем ущербе до 16 млн рублей на атаку; исследование Ростелекома и Института развития интернета - рост дипфейк-мошенничества в 2,3 раза за первое полугодие 2025 года при доле менее 1% общего объёма мошенничества. Приведённые значения иллюстрируют направление тенденции и не подлежат прямому сопоставлению.

2.7. Данные сетей верификации (уровни B/D)

Ключевое наблюдение: несколько независимых сетей с разными географиями и клиентскими базами сходятся в направлении тенденций - рост дипфейков, инъекций и синтетических личностей при снижении доли примитивных атак, - но существенно расходятся в абсолютных значениях вследствие различий определений и портфелей. Все данные раздела получены от заинтересованных сторон.

- **Entrust** (7-й ежегодный отчёт, ноябрь 2025 г.; более 1 млрд верификаций в 195 странах, период сентябрь 2024 - сентябрь 2025): дипфейки составили примерно каждую пятую обнаруженную биометрическую атаку; попытки с дипфейк-селфи выросли на 58%; инъекционные атаки - на 40% год к году; сложные цифровые подделки составили 35% документного мошенничества (против среднего значения

около 29% в 2022–2024 гг.); в секторах с бонусами за регистрацию (криптоактивы) до 67% мошенничества приходится на этап онбординга; в платёжном секторе 82% атак нацелены на процесс аутентификации. Поставщик оговаривает, что данные отражают его собственную сеть.

- **Sumsub** (ноябрь 2025 г.; проанализировано более 4 млн попыток мошенничества): глобальная доля мошенничества снизилась до 2,2% верификаций (с 2,6% в 2024 г.), однако сложные многоступенчатые атаки выросли на 180% год к году, а дипфейк-атаки - на 2100% за отчётный период (значение критически зависит от низкой базы и определения); рост синтетических документов +311% (I кв. 2024 - I кв. 2025). Прогноз резкого роста «агентных» атак в 2026 году является позицией поставщика.
- **LexisNexis Risk Solutions** (апрель 2026 г.; анализ 116 млрд транзакций): глобальный рост атак +8%; синтетические личности - 11% мирового мошенничества с восьмикратным ростом год к году; доля в Латинской Америке - 48,3% (максимум в мире); трафик, классифицированный поставщиком как «агентный», вырос на 450% - это ранний индикатор изменения автоматизации, но не прямая оценка объёма мошенничества, совершаемого ИИ-агентами.
- **iProov** (февраль 2025 г. и последующие данные): рост использования нативных виртуальных камер на 2665%; замен лица (face swap) - на 300% к 2023 году; около 24 000 участников в экосистеме «преступление-как-услуга»; инъекции против iOS выросли на 1151% во втором полугодии 2025 года (+741% за год).
- **Smile ID** (2026 г.; более 200 млн проверок, 35+ стран Африки): более 160 000 мошеннических попыток за один месяц сведены к 100 лицам, отдельные лица фигурировали более 12 000 раз на нескольких платформах; свыше 100 000 инъекционных попыток в месяц; в Южной Африке 87% отклонённых биометрических верификаций связаны с ИИ-имперсонацией и спуфингом; 90% заблокированного мошенничества выявлено по сигналам мобильного SDK (68% годом ранее). Данные - прямой независимый аналог «сетевого» вывода Unico о повторном использовании одних и тех же личностей против множества организаций.
- **Group-IB** (уровень D): свыше 8 000 дипфейк-попыток против одного банка за восемь месяцев; средний ущерб от дипфейк-вишинга в финансовом секторе - порядка 600 тыс. долларов на инцидент, более 10% опрошенных банков сообщали об отдельных потерях свыше 1 млн долларов.
- **Shufti** (прогноз, уровень D): проекция роста дипфейк-мошенничества на 495% в 2026 году и документных дипфейков на 3892% построена на аннуализации данных за январь–май 2026 года; сам поставщик маркирует значения как оценку (2026E). Используется исключительно как прогнозная гипотеза.

2.8. Экономика атакующего

По данным Group-IB и материалам ИНТЕРПОЛА, на криминальных площадках наблюдаются предложения дипфейк-услуг по минимальным ценам от 5 долларов, пакетов синтетических личностей - от 15 долларов, генерации дипфейк-изображений - 10–50 долларов, аренды продвинутого программного обеспечения замены лица - до 10 000 долларов. Указанные значения представляют собой **минимальные цены предложений, наблюдавшихся исследователями на криминальных площадках**, а не стоимость

успешной атаки: цена объявления не подтверждает качество услуги, факт сделки, полноту комплекта и способность пройти промышленную проверку. Тем не менее сходимость независимых наблюдений (Group-IB, ИНТЕРПОЛ, deepidv) позволяет с высокой уверенностью констатировать радикальное снижение входного барьера, качественно согласующееся с тезисом Unico о падении стоимости сложной атаки более чем на два порядка, хотя сам множитель «более 100x» остаётся внутренней оценкой Unico.

2.9. Человеческий фактор

В эксперименте iProov (пресс-релиз от 12 февраля 2025 г.; 2 000 респондентов из Великобритании и США) только **0,1% участников безошибочно классифицировали весь предложенный набор** реальных и синтетических материалов, при этом 57% выражали уверенность в своей способности отличать подделки, а свыше 60% сохраняли уверенность независимо от правильности ответов. Результат демонстрирует ограниченность невооружённой человеческой проверки и избыточную уверенность пользователей, однако **не является оценкой точности распознавания отдельного дипфейк-объекта**: вероятность безошибочно пройти серию испытаний быстро снижается с ростом числа стимулов, а полная методология исследования (число предъявлений, доля подделок, чувствительность и специфичность по объектам) поставщиком не раскрыта. Смежные данные согласуются по направлению: по данным deepidv, операторы ручной проверки верно распознавали качественные дипфейки примерно в четверти случаев. Совокупный вывод: ручная визуальная проверка не может выступать основным контролем против современных синтетических подделок; человек в контуре сохраняет ценность для расследования и оспаривания, но не для первичной детекции.

2.10. Показательный инцидент: дело Arup (уровень A4, рамка отдельного случая)

Январь–февраль 2024 года, Гонконг: сотрудник финансовой службы инженерной компании Arup провёл 15 переводов на сумму около 200 млн гонконгских долларов (порядка 25,6 млн долларов США) после видеоконференции, в которой финансовый директор и коллеги оказались дипфейками; инцидент подтверждён полицией Гонконга, компания публично признала себя потерпевшей в мае 2024 года. Случай используется в исследовании как доказательство возможного масштаба одного успешного инцидента и уязвимости процедур подтверждения распоряжений, но не как статистическое подтверждение распространённости подобных атак.

3. Часть II. Технологии и практики противодействия

3.1. Пятислойная модель защиты

Современный контур противодействия строится из пяти взаимодополняющих слоёв. Ни один слой в отдельности не обеспечивает достаточной защиты: качественный алгоритм проверки живости бесполезен, если злоумышленник подменяет видеопоток между камерой и программным модулем; надёжная первичная проверка бесполезна, если восстановление доступа защищено слабее регистрации.

1. **Сенсор и доверенный захват:** активная и пассивная проверка живости (liveness detection); детекция презентационных атак (PAD, стандарт ISO/IEC 30107-3); детекция инъекций и виртуальных камер; контроль целостности видеопотока и конвейера захвата; аппаратная аттестация устройства и приложения (hardware attestation; Google Play Integrity; Apple App Attest); доверенный захват фотографии для документов (trusted capture).
2. **Аналитика:** детекция дипфейков (видео, голос, документы); выявление синтетических личностей; интеллект по устройствам и цифровые отпечатки (device intelligence, fingerprinting); поведенческая биометрия как вероятностный сигнал риска, а не самостоятельное доказательство личности; графовый анализ связей между документами, лицами, устройствами, счетами и получателями платежей.
3. **Процесс:** непрерывная оценка риска после первичной проверки; защита процедуры восстановления доступа на уровне не ниже первичной регистрации; усиленная (step-up) аутентификация перед критическими действиями; подписание операций; защита сессии.
4. **Сетевой уровень:** межорганизационный обмен сигналами о мошенничестве (консорциумы, национальные базы наподобие Cifas); архитектуры приватного сопоставления (пересечение множеств с сохранением конфиденциальности, федеративное обучение, доверенные среды исполнения); обязательные механизмы отзыва ошибочного флага, оспаривания и журналирования обращений - сеть, распространяющая ошибку, распространяет и вред.
5. **Криптографическая инфраструктура доверия:** проверяемые учётные данные (W3C Verifiable Credentials 2.0); европейский кошелек цифровой идентичности (eIDAS 2.0); мобильные удостоверения (ISO/IEC 18013-5); фишинго-устойчивая аутентификация (passkeys, FIDO2/WebAuthn); криптографическая проверка документов, включая чтение NFC-чипов; происхождение контента (C2PA).

Практически значимый сдвиг, подтверждаемый данными Smile ID: до 90% блокируемого мошенничества выявляется по сигналам мобильного SDK и устройства, а не по анализу изображения, - центр тяжести защиты смещается от «смотреть на картинку» к «доверять каналу захвата».

3.2. Независимые государственные испытания DHS S&T RIVTD (уровень A3, по первичным документам)

Демонстрация технологий удалённой проверки личности (Remote Identity Validation Technology Demonstration, RIVTD) проведена Директоратом науки и технологий

Министерства внутренней безопасности США (DHS S&T) в партнёрстве с TSA, криминалистической лабораторией HSI и NIST на площадке Maryland Test Facility (Аппер-Марлборо, Мэриленд) в 2023–2024 годах. Согласно официальной странице прошедших событий MdTF, продемонстрированы: **9 систем валидации документов** (трек 1), **16 систем сопоставления селфи с документом** (трек 2), **6 активных и 15 пассивных подсистем детекции презентационных атак** (трек 3). В официальных публикациях участники **анонимизированы** (обозначения MTDS 1–16, PAD-A 1–6, PAD-P 1–15); соответствие псевдонимов конкретным поставщикам может быть запрошено через DHS. Настоящий раздел опирается исключительно на первичные публикации MdTF/DHS; матрица первичных документов приведена в таблице ниже, полные официальные таблицы результатов - в Приложении Б.

Документ	Тре к	Статус	Адрес
MdTF. RIVTD Results (таблицы результатов треков 2–3; инфографика по всем трекам; обновление февраль 2025 г.)	1–3	Официальн ая страница результатов	https://mdtf.org/rivtd/Results2023
MdTF. Past Events: RIVTD 2023 (состав продемонстрирован ных систем)	1–3	Официальн ая страница	https://mdtf.org/rivr/PastEvents
MdTF. 2023 RIVTD Technical Clarification Webinar (методика трека 1 по NIST SP 800-63A)	1	Официальн ая презентаци я (PDF)	https://mdtf.org/Downloads/2023_RIVTD_TechnicalClarificationWebinar.pdf
MdTF. RIVTD Document Validation FAQ	1	Официальн ый документ (PDF)	https://mdtf.org/Downloads/RIVTD Document Validation FAQ.pdf
DHS S&T. Программная страница RIVTD (описание треков; переход к программе-преемнику RIVR)	1–3	Официальн ая страница	https://www.dhs.gov/science-and-technology/remote-identity-validation-technology-demonstration

Методология по трекам (по первичным документам)

Трек	Задача и выборка	Метрики	Ограничения
1	Валидация американских водительских удостоверений по изображениям лицевой и обратной сторон; методика согласно NIST SP 800-63A; 9 систем	Табличные метрики в открытых первичных публикациях отсутствуют; результаты представлены инфографикой	Открытые количественные выводы уровня A3 по треку 1 недоступны; встречающиеся в отраслевой прессе детали (состав выборки, DFAR/DFRR, рекомендованный порог ошибок) имеют уровень E

Трек	Задача и выборка	Метрики	Ограничения
2	Сопоставление селфи с фотографией документа; 16 систем; изображения 1 633 оплачиваемых добровольцев (сборы в Мэриленде и Калифорнии)	Доля неизвлечения селфи и документа (FTXR, отдельно); доля ложного несовпадения (FNMR) при порогах поставщика с целевым FMR 1:10 000; проверка калибровки порогов на многолетнем эталонном наборе MdTF (классификация: консервативный / пермиссивный / ожидаемый)	Часть систем не оценена (NA) из-за технических сбоев; пороги задавались поставщиками
3	Детекция презентационных атак; 6 активных и 15 пассивных подсистем; 660 участников; несколько моделей смартфонов и несколько видов атак	BPCER и APCER; отчётные значения - худший случай (максимум) по смартфонам и видам атак; для активных дополнительно среднее время транзакции и доля удовлетворённых пользователей, для пассивных - среднее время обработки	Опубликованный контур охватывает презентационные атаки; цифровые инъекции и виртуальные камеры в него не входят

Официальные результаты (подтверждённые первичными таблицами)

- **Трек 2.** Слабым звеном оказалось извлечение данных из документа: доля неизвлечения документа варьировала от 0% до «<100%» (две системы практически не справились с задачей), тогда как доля неизвлечения селфи не превышала 9%. Доля ложного несовпадения составила от 0% до «<100%»; у четырёх систем метрика не рассчитана из-за технических сбоев. Калибровка порогов подтверждена как самостоятельная проблема: пермиссивные пороги повышают риск совпадения самозванца с чужим документом, консервативные - риск отклонения законного владельца.
- **Трек 3.** Лучшие пассивные подсистемы сочетали в худшем сценарии ошибку на добросовестных предъявлениях (BPCER) не выше 0,5% с ошибкой на атаках (APCER) не выше 3,3%; при этом у четырёх пассивных подсистем максимальный APCER составил 96,7–100%, то есть в худшем случае атаки отдельных видов пропускались почти полностью. Активные подсистемы достигали APCER 0%, но при BPCER (максимум) от 6,1% до 58,6% и среднем времени транзакции 27–40 секунд - количественное подтверждение компромисса между устойчивостью к атакам и удобством легитимных пользователей.
- **Общий вывод.** Первичные данные документируют существенный разброс качества между участниками в пределах каждого трека и подтверждают, что успешное выполнение одного компонента удалённой проверки личности не гарантирует надёжности процесса в целом.

Анонимизация и атрибуция участников

Поскольку официальные публикации анонимизируют участников, установление соответствия «псевдоним - поставщик - метрика - порог - дата» по открытым первичным документам невозможно. Отдельные компании самостоятельно заявили о своём участии и результатах в собственных пресс-релизах; такие заявления имеют уровень доказательности D (заинтересованная сторона) и в выводы уровня A3 настоящего исследования не

переносятся. Термины «лидер», «победитель» и «лучший» не используются, поскольку DHS ранжирования участников не устанавливает. Программой-преемником с 2025 года является Remote Identity Validation Rally (RIVR) с ужесточёнными целевыми порогами; её результаты публикуются MdTF в том же анонимизированном формате.

3.3. Прочие независимые испытания (уровень A3)

NIST FATE-PAD (отчёт NISTIR 8491; протестировано 82 алгоритма) предоставляет независимую методологическую и экспериментальную основу для оценки **пассивных программных** алгоритмов детекции презентационных атак, работающих с обычными двумерными изображениями. Направление в настоящее время закрыто для приёма новых заявок; опубликованные результаты не охватывают современный контур инъекционных атак, виртуальных камер, аппаратной аттестации и дипфейков в реальном времени и не должны представляться как непрерывно обновляемый рейтинг продуктов 2026 года.

NIST FATE MORPH и руководство по морфингу (2025 г.): риск морфинга особенно высок, когда заявитель самостоятельно предоставляет фотографию с неизвестной историей происхождения; предпочтительный контроль - доверенный захват; при его невозможности - автоматическая детекция морфинга с обязательным последующим расследованием срабатываний с учётом стоимости ложных решений.

Общий вывод раздела: заявления поставщиков о точности («99,3%» и аналогичные) без раскрытия знаменателя, операционной точки, состава атак и независимого подтверждения (DHS, NIST, iBeta по ISO/IEC 30107-3) не подлежат использованию в сравнительных выводах и закупочных решениях.

3.4. Рынок решений: доказательная матрица вместо перечня

Международный рынок представлен, в частности, компаниями Entrust/Onfido, Sumsb, iProov, FaceTec, Jumio, Veriff, Incode, Persona, Socure, Mitek, Trulioo, BioCatch, Pindrop (голос), Reality Defender, ROC.ai, Paravision, IDEMIA, Aware; российский - решениями VisionLabs (детекция дипфейков и проверка живости, сертификация iBeta уровней 1/2; программно-аппаратный комплекс «ВЛ-Био» класса защиты KC3), NtechLab (FindFace, автономный сервис проверки живости), ЦРТ (голосовая биометрия и антиспуфинг), IDX, а также инфраструктурой Единой биометрической системы. Перечисление не является рейтингом. Заявления отдельных поставщиков об «успешном прохождении» испытаний RIVTD представляют собой самоидентификацию в анонимизированных государственных испытаниях и имеют уровень доказательности D (см. раздел 3.2).

Для объективного сравнения каждому решению должна заполняться единая доказательная карточка; включение продукта в сравнительный вывод допустимо только при наличии проверяемой технической информации и независимых испытаний.

Параметр карточки	Содержание
Защищаемый этап	Регистрация; аутентификация; восстановление; операция
Класс атак	Презентационные; инъекции; виртуальные камеры; морфинг; клон голоса; синтетическая личность
Тип решения	API; SDK; облачный сервис; локальное развёртывание

Параметр карточки	Содержание
Независимые испытания	DHS; NIST; iBeta (ISO/IEC 30107-3); иные лаборатории - с датой, треком и версией
Метрики	APCER; BPCER; FAR; FRR; задержка; операционная точка
Тестовый набор	Размер; типы атак; период; разделение обучающей и тестовой выборки
Аттестация устройства	Android; iOS; веб-канал
Инъекционные атаки	Какие именно векторы обнаруживаются
Объяснимость и аудит	Доступные заказчику причины решения; журналирование; версия модели
Конфиденциальность	Состав собираемых данных; хранение; трансграничная передача; сроки
Статус доказательства	A3 / B / C / D
Ограничения	Где решение не проверено или неприменимо

3.5. Регуляторный контекст (уровень A2, сверено с первичными актами)

Настоящий раздел представляет аналитический обзор регулирования и не является индивидуальным юридическим заключением. Утверждения сверены с первичными и консолидированными текстами актов; аналитические интерпретации авторов помечены отдельно. Полная верификационная таблица юридически значимых утверждений и матрица применимости приведены в Приложении В.

Регламент ЕС об искусственном интеллекте (Регламент (ЕС) 2024/1689; вступил в силу 1 августа 2024 г.). График применения установлен статьёй 113: запрещённые практики (статья 5) - со 2 февраля 2025 г.; обязательства для моделей ИИ общего назначения - со 2 августа 2025 г.; общая применимость - со 2 августа 2026 г. Регламент о внесении изменений («Цифровой омнибус по ИИ»: предложение Комиссии от 19 ноября 2025 г.; политическое согласие Совета и Парламента 7 мая 2026 г.; одобрение Европарламента 16 июня 2026 г., 423 голоса против 57; окончательное одобрение Совета 29 июня 2026 г.; вступил в силу 27 июля 2026 г.) изложил статью 113 в новой редакции: обязательства для высокорисковых систем по статье 6(2) и приложению III (включая биометрию) применяются со **2 декабря 2027 г.**, для систем по статье 6(1) и приложению I (встроенных в регулируемую продукцию) - со 2 августа 2028 г.

Границы биометрического основания установлены самим приложением III: пункт 1(a) **прямо исключает** из высокорисковой категории системы биометрической верификации, единственной целью которых является подтверждение того, что конкретное физическое лицо является тем, за кого себя выдаёт (см. также определения удалённой биометрической идентификации и верификации в статье 3). Это исключение не распространяется автоматически на системы подтверждения личности в целом, оценки риска и оркестрации противодействия мошенничеству. Пункт 5(b) приложения III относит к высокорисковым системы оценки кредитоспособности и кредитного скоринга физических лиц, **за исключением систем, используемых для выявления финансового мошенничества**; данное изъятие привязано именно к контексту кредитоспособности и не является универсальным для любого антифрод-решения.

Аналитическая интерпретация (позиция авторов, не установленная норма). Классификация комплексной антифрод-системы зависит от её назначения, выходных решений и фактического использования; по мере того как сигнал о мошенничестве

начинает определять отказ в кредите или ограничение обслуживания, применимость изъятия становится предметом толкования. Данная позиция опирается на подготавливаемые Комиссией руководства по применению статьи 6(3) и совпадает с аналитической оценкой Liminal; правоприменительная практика по этому вопросу ещё не сложилась.

eIDAS 2.0 (Регламент (ЕС) 2024/1183; вступил в силу 20 мая 2024 г.): государства-члены обязаны предоставить гражданам и резидентам европейские кошельки цифровой идентичности в течение 24 месяцев после принятия имплементационных актов; акты приняты в конце 2024 года, что определяет срок предоставления кошельков к концу 2026 года. Кошелёк снижает часть рисков подделки документов, но переносит угрозы на уровень устройства, процедуры выдачи и злоупотреблений проверяющей стороны (компрометация кошелька, избыточное раскрытие атрибутов, недобросовестный запрашивающий).

NIST SP 800-63, четвёртая редакция (финальная версия 31 июля 2025 г., итог процесса с двумя публичными проектами и почти 6 000 комментариев): риск-ориентированная модель управления цифровой идентичностью; три уровня уверенности (подтверждение личности, аутентификация, федерация); прямое признание passkeys и цифровых кошельков; курс на фишинго-устойчивую аутентификацию; рекомендованные метрики непрерывной оценки.

W3C Verifiable Credentials 2.0 (рекомендация от 15 мая 2025 г.): криптографическая целостность, выборочное раскрытие, механизм статуса и отзыва. Принципиальная оговорка самого стандарта: успешная криптографическая проверка удостоверения подтверждает его происхождение и целостность, но **сама по себе не доказывает истинность содержащихся в нём утверждений** - качество экосистемы определяется процедурами выдачи.

Российская Федерация. Идентификация и аутентификация с использованием биометрических персональных данных регулируется Федеральным законом от 29.12.2022 № 572-ФЗ во взаимосвязи со 152-ФЗ. Закон устанавливает **закрытый перечень видов используемой биометрии - изображение лица и запись голоса**. Исходные биометрические образцы размещаются в государственной информационной системе «Единая биометрическая система»; сбор осуществляют аккредитованные организации (банки, МФЦ и иные уполномоченные субъекты). Организации, проводящие идентификацию и аутентификацию, работают не с исходными образцами, а с **векторами** - производными математическими представлениями; хранение исходных изображений и записей в собственных системах коммерческих организаций вне установленных законом случаев не допускается. Векторы имеют особый правовой режим: в силу части 5 статьи 3 572-ФЗ к ним не применяются положения статьи 11 и меры обеспечения безопасности, установленные в соответствии со статьёй 19 Федерального закона № 152-ФЗ.

Категории регистрации биометрии, употребляемые на практике («упрощённая», «стандартная» - дистанционные способы через приложения, «подтверждённая» - при личном присутствии в банке или МФЦ), закреплены на подзаконном и операторском уровне; сам закон № 572-ФЗ эти наименования не вводит, что необходимо учитывать при цитировании. Федеральным законом от 01.04.2025 № 41-ФЗ в статью 7 Федерального закона № 115-ФЗ введён пункт 5.8-2: микрофинансовые организации при дистанционном приёме на обслуживание физических лиц обязаны проводить идентификацию с использованием ЕСИА и Единой биометрической системы, а при дистанционном

заключении каждого договора займа - аутентификацию клиента через ЕСИА и ЕБС. Для микрофинансовых компаний обязанность применяется с **1 марта 2026 года**, для микрокредитных компаний - с 1 марта 2027 года; согласно информации Банка России от 5 мая 2026 года, до 1 января 2027 года надзорные меры к микрофинансовым компаниям за неисполнение этой обязанности не применяются. По разъяснениям участников рынка, для указанных целей используется подтверждённая биометрия (уровень Е - практика, а не буквальная норма). Законодательные инициативы о признании использования дипфейков отягчающим обстоятельством и об обязательной маркировке контента, созданного искусственным интеллектом, по состоянию на август 2026 года имеют статус **внесённых законопроектов** и не являются действующими нормами.

3.6. Тренды на горизонте 12–24 месяцев

- Переход от разовой проверки к непрерывной верификации личности и риска на протяжении всей сессии и жизни учётной записи.
- «Агентное» мошенничество: полностью автоматизированные кампании, о которых предупреждают ИНТЕРПОЛ, Sumsb и LexisNexis; развитие детекции нечеловеческих паттернов взаимодействия на уровне сессии; при этом рост «агентного» трафика сам по себе не равен росту мошенничества, совершаемого агентами, и требует отдельного учёта.
- Происхождение контента: стандарты C2PA и криптографические водяные знаки как дополнение (но не замена) детекции подделок.
- Стандартизация детекции инъекционных атак (в том числе европейская техническая спецификация CEN/TS 18099) и смещение независимых испытаний от чисто презентационных атак к цифровым векторам - актуальная задача с учётом того, что опубликованный контур RIVTD цифровые инъекции не охватывал.
- Инфраструктурный сдвиг к криптографически проверяемым удостоверениям (кошельки eIDAS 2.0, mDL, verifiable credentials) с одновременным переносом рисков на процедуры выдачи и восстановления.

4. Часть III. Научный задел

4.1. Организация по исследовательским вопросам

Научный ландшафт систематизирован не по публикациям, а по двенадцати нерешённым исследовательским вопросам, определяющим практическую применимость технологий.

1. Переносимость между наборами данных (cross-dataset generalisation): работает ли детектор против генератора, отсутствовавшего при обучении.
2. Детекция в открытом множестве (open-set): атаки, не представленные при обучении.
3. Детекция инъекций в реальном времени, включая виртуальные камеры и подмену потока.
4. Единая детекция физических и цифровых атак (unified physical–digital attack detection).
5. Мультимодальная аудиовизуальная детекция.
6. Устойчивость к состязательным воздействиям (adversarial robustness), включая атаки на детекторы аудиодипфейков.
7. Детекция морфинга (одиночное изображение и дифференциальная).
8. Графовое выявление синтетических личностей и мошеннических колец: гетерогенные и временные графовые нейронные сети, устойчивость к отравлению графа, объяснимость, масштабирование.
9. Коллаборативная детекция с сохранением конфиденциальности: федеративное обучение, приватное пересечение множеств, многосторонние вычисления, гомоморфное шифрование, доверенные среды исполнения - с обязательным разграничением демонстрационных экспериментов и производственной инфраструктуры.
10. Справедливость и демографические различия качества.
11. Калиброванная неопределённость и воздержание от решения (abstention).
12. Совместная работа человека и машины при проверке и оспаривании.

4.2. Ключевые публикации и эталонные наборы

- Yu Z., Qin Y., Li X., Zhao C., Lei Z., Zhao G. Deep Learning for Face Anti-Spoofing: A Survey // IEEE Transactions on Pattern Analysis and Machine Intelligence. 2023. Т. 45, № 5. С. 5609–5631 (онлайн 19.10.2022). DOI: 10.1109/TPAMI.2022.3215850. Канонический обзор глубинных методов детекции спуфинга лица; сопровождается актуализируемым перечнем литературы.
- Tolosana R., Vera-Rodriguez R., Fierrez J., Morales A., Ortega-Garcia J. DeepFakes and Beyond: A Survey of Face Manipulation and Fake Detection // Information Fusion. 2020. Т. 64. С. 131–148. DOI: 10.1016/j.inffus.2020.06.014. Классификация четырёх типов манипуляций лицом.

- A Survey on Speech Deepfake Detection. arXiv:2404.13914 (препринт; редакция июля 2025 г.): детекция речевых дипфейков, межнаборная оценка, защита от состязательных атак.
- Principles of Designing Robust Remote Face Anti-Spoofing Systems. arXiv:2406.03684 (препринт): экспериментальная демонстрация слабой переносимости современных моделей на цифровые инъекционные атаки - состязательный шум, реалистичные дипфейки, цифровой повтор.
- Deepfake-Eval-2024. arXiv:2503.02857 (препринт): на «диких» дипфейках 2024 года площадь под кривой у открытых детекторов падает на 45–50% относительно академических эталонов - прямое количественное подтверждение проблемы переносимости.
- Обзоры применения графовых нейронных сетей к финансовому мошенничеству (в т. ч. arXiv:2411.05815, препринт): подтверждённая применимость при открытых проблемах масштабирования, объяснимости, временной динамики и переноса моделей между организациями.
- Эталонные наборы и конкурсы: FaceForensics++; Celeb-DF; DFDC; ASVspoof (челлендж); SiW-Mv2; UniAttackData; профильные секции CVPR/ICCV. Высокие результаты на академических наборах не переносятся автоматически на устойчивость к «живой» инъекции в мобильном приложении.

4.3. Сводный вывод по научному разделу

Научное состояние области характеризуется зрелостью методов против известных, статичных, презентационных атак и систематическим отставанием против неизвестных генераторов, цифровых инъекций и состязательных воздействий. Для практики это означает: закупочные решения не должны опираться на результаты, полученные на академических эталонах; обязательны испытания на неизвестных для модели атаках, оценка деградации при сжатии и повторной съёмке, отдельная проверка презентационного и инъекционного контуров, контроль демографической однородности качества.

5. Рекомендации для организаций

5.1. Триггеры приоритизации от собственных данных

Рекомендации отвязаны от спорных отраслевых процентов (доли атак в сетях отдельных поставщиков зависят от их клиентской базы и не являются универсальными нормативами). Приоритизация строится от собственных операционных данных организации.

Критический приоритет (немедленное реагирование)

- Подтверждён успешный инъекционный обход проверки.
- Одна сущность (лицо, документ, устройство) проходит проверки в нескольких юридических лицах группы.
- Отсутствует механизм отзыва ошибочного сетевого флага при участии в обмене сигналами.
- Биометрическое решение влияет на кредит, счёт или доступ к значимой услуге без объяснения причин и процедуры апелляции.

Немедленный приоритет

- Зафиксированы инциденты с инъекциями или виртуальными камерами.
- Мобильное приложение не использует аппаратную аттестацию (Play Integrity, App Attest).
- Процедура восстановления доступа защищена слабее первичной регистрации.
- Отсутствует анализ риска на уровне сессии.
- Ручная визуальная проверка является основным контролем против дипфейков.
- Решения поставщика невозможно аудировать (нет причин решения, версии модели, журналирования).

Повышенный приоритет

- Растёт доля повторно используемых лиц, устройств или документов.
- Один идентификатор (телефон, почта, устройство) связан с множеством учётных записей.
- Увеличивается число атак после успешной первичной проверки (захваты, подмены SIM-карт).
- Наблюдается расхождение сигналов проверки документов и интеллекта по устройствам.
- Растут ложные срабатывания или очередь ручной проверки.

5.2. Этапный план

Этап 1 (0–3 месяца). Внедрить детекцию инъекций и виртуальных камер совместно с аттестацией устройства; исключить ручную проверку как единственный контроль против дипфейков; запрашивать у поставщиков результаты независимых испытаний (DHS, NIST, iBeta по ISO/IEC 30107-3) с датой, треком, составом атак и операционной точкой вместо

самоотчётных метрик; с учётом худших сценариев RIVTD (APCER до 100% у отдельных пассивных подсистем) требовать раскрытия максимальных, а не средних значений ошибок по видам атак.

Этап 2 (3–12 месяцев). Построить многослойный контур (проверка живости, детекция инъекций, анализ генеративных артефактов, поведенческие и аппаратные сигналы); порог допуска определять по результатам независимого испытания на репрезентативной для организации выборке, измеряя одновременно долю успешных атак, APCER, BPCER, FAR, FRR, долю ручной проверки и операционные потери от ложных решений; перейти к непрерывной оценке риска; выровнять защиту восстановления доступа с первичной регистрацией; внедрить графовый анализ связей; при участии в консорциумах обмена сигналами - только при наличии приватного сопоставления, срока действия флагов, процедуры апелляции и механизма отзыва ошибочной отметки во всех участвующих системах.

Этап 3 (12–24 месяца). Подготовка к применению обязательств Регламента ЕС об ИИ (2 декабря 2027 г.) и к экосистеме eIDAS 2.0; внедрение проверяемых учётных данных, мобильных удостоверений, passkeys/FIDO2 и происхождения контента; построение защиты от «агентных» атак через поведенческую аналитику сессий; полный контур управления решениями - журналирование причин, версий моделей и источников сигналов, независимая проверка, апелляция, мониторинг ложных приёмов и отказов, контроль демографической справедливости, регулярные учения «красной команды».

5.3. Обязательные вопросы поставщику до пилота или закупки

1. Полная методология фигурирующих в маркетинге показателей: числитель, знаменатель, период, география, состав выборки.
2. Результаты независимых испытаний с указанием лаборатории, трека, версии продукта и даты; для проверки живости - соответствие ISO/IEC 30107-3; для участников RIVTD/RIVR - официальное подтверждение соответствия псевдонима, полученное через DHS.
3. Раздельные показатели по презентационным и инъекционным атакам; перечень обнаруживаемых векторов инъекций и виртуальных камер.
4. APCER, BPCER, FAR, FRR в заявленной операционной точке - с указанием максимальных значений по видам атак и устройствам, а не только средних; показатели по регионам, устройствам и демографическим группам.
5. Механизм объяснения решения, журналирование, версия модели, порядок аудита.
6. Процедура исправления ошибочного флага, в том числе распространённого по сети; сроки; ответственность участников.
7. Состав собираемых данных; сроки хранения биометрии, изображений и производных признаков; регионы размещения; трансграничная передача; субподрядчики.
8. Соответствие применимому регулированию: Регламент ЕС об ИИ, GDPR, 152-ФЗ/572-ФЗ, отраслевые требования.

6. Ограничения исследования

- **Конфликт интересов.** Данные Entrust, Sumsb, iProov, Smile ID, Group-IB, Shufti и deepidv получены от поставщиков защитных решений. Направления тенденций подтверждаются сходимостью независимых сетей; абсолютные значения не сопоставимы и не суммируемы.
- **Несопоставимость метрик.** Попытки, успешные атаки, экспозиция и фактические потери - разные величины; проценты роста зависят от базы и определения. Все ключевые показатели снабжены паспортами (раздел 1.3).
- **Оценочный характер глобальных сумм.** Значение 442 млрд долларов (ИНТЕРПОЛ) - институциональная оценка порядка величины; цифра «свыше 400 млрд долларов» из отчёта Liminal + Unico первичным источником не подтверждена и трактуется как агрегат с нераскрытой методикой.
- **Прогнозы - гипотезы.** Проекция Shufti (+495% и +3892% на 2026 г.), оценка Deloitte (около 40 млрд долларов потерь в США к 2027 г. при консервативном сценарии около 22 млрд) и прогноз роста потерь финансовых организаций на 121% к 2030 году (Liminal) - экстраполяции, используемые только для планирования.
- **Серийный характер теста iProov.** Показатель 0,1% относится к безошибочному прохождению всей серии стимулов и не является точностью распознавания отдельного объекта; полная методика закрыта.
- **Вендорские пороги.** Показатель deepidv «обход менее 0,3% при многослойной защите» - иллюстративный результат заинтересованной стороны на собственной выборке, а не отраслевой норматив.
- **Российский контур.** Официальная статистика Банка России не выделяет дипфейки; отраслевые оценки (Сбербанк, «Информзащита», Ростелеком/ИРИ) основаны на разных методиках; открытой статистики попыток атак на саму ЕБС не обнаружено.
- **Государственные испытания RIVTD.** Раздел 3.2 сверен с первичными публикациями MdTF/DHS; остающиеся ограничения самих первичных данных: анонимизация участников исключает открытую атрибуцию результатов конкретным поставщикам; отчётные метрики трека 3 представляют худший случай по устройствам и видам атак; опубликованный контур не охватывает цифровые инъекции; по треку 1 открытые табличные результаты отсутствуют. Результаты NIST FATE-PAD ограничены пассивными программными алгоритмами и закрытым набором заявок.
- **Нормативный раздел.** Раздел 3.5 и Приложение В представляют аналитический обзор регулирования, сверенный с первичными актами, а не индивидуальное юридическое заключение; аналитические интерпретации (в частности, об объёме изъятия для выявления мошенничества по приложению III Регламента ЕС об ИИ) помечены как позиция авторов.
- **Исключённые показатели.** Утверждение об атрибуции 2% поддельных документов конкретным генеративным моделям исключено из числовых выводов до раскрытия методики атрибуции.

7. Реестр источников

Уровень А1. Официальная статистика и административные данные

- FBI, Internet Crime Complaint Center (IC3). 2025 Internet Crime Report. Апрель 2026. URL: https://www.ic3.gov/AnnualReport/Reports/2025_IC3Report.pdf (дата обращения: 05.08.2026).
- INTERPOL. Global Financial Fraud Threat Assessment. 2-е изд. Март 2026. URL: <https://www.interpol.int/News-and-Events/News/2026/INTERPOL-report-warns-of-increasingly-sophisticated-global-financial-fraud-threat> (дата обращения: 05.08.2026).
- UK Finance. Annual Fraud Report 2026. Июнь 2026. URL: <https://www.ukfinance.org.uk/policy-and-guidance/reports/annual-fraud-report-2026> (дата обращения: 05.08.2026).
- Cifas. Fraudscape 2026. Март 2026. URL: <https://www.cifas.org.uk/newsroom/fraudscape2026> (дата обращения: 05.08.2026).
- European Banking Authority; European Central Bank. Joint Report on Payment Fraud. URL: <https://www.eba.europa.eu/publications-and-media/press-releases/joint-eba-ecb-report-payment-fraud-strong-authentication-remains-effective-fraudsters-are-adapting> (дата обращения: 05.08.2026).
- Банк России. Обзор операций, совершённых без добровольного согласия клиентов финансовых организаций: за 2024 год; за 2025 год. URL: https://www.cbr.ru/analytics/ib/operations_survey/2024/; https://www.cbr.ru/analytics/ib/operations_survey/2025/ (дата обращения: 05.08.2026).

Уровень А2. Нормативные акты и официальные руководства

- Регламент (ЕС) 2024/1689 Европейского парламента и Совета (Регламент об искусственном интеллекте). URL: https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ%3AL_202401689 ; сроки применения: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai> (дата обращения: 05.08.2026).
- Регламент (ЕС) 2024/1183 (eIDAS 2.0). Вступил в силу 20.05.2024. URL: <https://eur-lex.europa.eu/eli/reg/2024/1183/oj> (дата обращения: 05.08.2026).
- Регламент о внесении изменений в Регламент (ЕС) 2024/1689 («Цифровой омнибус по ИИ»). Одобрен Европарламентом 16.06.2026 и Советом 29.06.2026; вступил в силу 27.07.2026. Хронология и итоговый текст: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai> (дата обращения: 05.08.2026).
- Федеральный закон от 01.04.2025 № 41-ФЗ (пункт 5.8-2 статьи 7 Федерального закона № 115-ФЗ: идентификация и аутентификация клиентов МФО через ЕСИА и ЕБС); Информация Банка России от 05.05.2026 о неприменении надзорных мер до 01.01.2027 (дата обращения: 05.08.2026).

- European Commission. European Digital Identity (eIDAS 2.0). URL: https://commission.europa.eu/topics/digital-economy-and-society/european-digital-identity_en (дата обращения: 05.08.2026).
- NIST. Special Publication 800-63-4: Digital Identity Guidelines. Финальная версия 31.07.2025. DOI: 10.6028/NIST.SP.800-63-4. URL: <https://csrc.nist.gov/pubs/sp/800/63/4/final> ; <https://pages.nist.gov/800-63-4> (дата обращения: 05.08.2026).
- W3C. Verifiable Credentials 2.0 (семейство спецификаций). Рекомендация 15.05.2025. URL: <https://www.w3.org/news/2025/the-verifiable-credentials-2-0-family-of-specifications-is-now-a-w3c-recommendation/> (дата обращения: 05.08.2026).
- Федеральный закон от 29.12.2022 № 572-ФЗ; Федеральный закон от 27.07.2006 № 152-ФЗ; Банк России - материалы о категориях биометрии ЕБС. URL: <http://kremlin.ru/acts/bank/48740> ; https://www.cbr.ru/about_br/publ/results_work/2025/razvitie-tekhnologiy-i-podderzhka-innovaciy/ (дата обращения: 05.08.2026).

Уровень А3. Государственные технические испытания (первичные публикации)

- Maryland Test Facility. Remote Identity Validation Technology Demonstration Results (треки 2–3; инфографика по трекам 1–3). URL: <https://mdtf.org/rivtd/Results2023> (дата обращения: 05.08.2026).
- Maryland Test Facility. Past Events: RIVTD 2023 (состав продемонстрированных систем). URL: <https://mdtf.org/rivr/PastEvents> (дата обращения: 05.08.2026).
- Maryland Test Facility. 2023 RIVTD Technical Clarification Webinar (PDF). URL: https://mdtf.org/Downloads/2023_RIVTD_TechnicalClarificationWebinar.pdf (дата обращения: 05.08.2026).
- Maryland Test Facility. RIVTD Document Validation FAQ (PDF). URL: <https://mdtf.org/Downloads/RIVTD Document Validation FAQ.pdf> (дата обращения: 05.08.2026).
- DHS Science and Technology Directorate. Remote Identity Validation Technology Demonstration (программная страница). URL: <https://www.dhs.gov/science-and-technology/remote-identity-validation-technology-demonstration> (дата обращения: 05.08.2026).
- NIST. Face Analysis Technology Evaluation, Presentation Attack Detection (NISTIR 8491; страница направления). URL: https://pages.nist.gov/frvt/html/frvt_pad.html (дата обращения: 05.08.2026).
- NIST. FATE MORPH; практическое руководство по детекции морфинга (2025). URL: https://pages.nist.gov/frvt/reports/morph/frvt_morph_report.pdf (дата обращения: 05.08.2026).

Уровень А4. Расследованные инциденты

- Дело Arup, Гонконг, 2024: подтверждение полиции Гонконга; публичное признание компании (май 2024). Освещение: CNN, 16.05.2024. URL: <https://www.cnn.com/2024/05/16/tech/arup-deepfake-scam-loss-hong-kong-intl-hnk> (дата обращения: 05.08.2026).

Уровень В. Научные публикации (препринты помечены)

- Yu Z. et al. Deep Learning for Face Anti-Spoofing: A Survey // IEEE TPAMI. 2023. DOI: 10.1109/TPAMI.2022.3215850 (дата обращения: 05.08.2026).
- Tolosana R. et al. DeepFakes and Beyond // Information Fusion. 2020. DOI: 10.1016/j.inffus.2020.06.014 (дата обращения: 05.08.2026).
- A Survey on Speech Deepfake Detection. arXiv:2404.13914 (препринт, ред. 07.2025) (дата обращения: 05.08.2026).
- Principles of Designing Robust Remote Face Anti-Spoofing Systems. arXiv:2406.03684 (препринт) (дата обращения: 05.08.2026).
- Deepfake-Eval-2024. arXiv:2503.02857 (препринт) (дата обращения: 05.08.2026).
- Graph Neural Networks for Financial Fraud Detection (обзор). arXiv:2411.05815 (препринт) (дата обращения: 05.08.2026).

Уровень С. Отраслевые исследования

- TransUnion. Synthetic Identity Fraud Analysis; H1 2025 Omnichannel Fraud Update. URL: <https://newsroom.transunion.com/transunion-research-highlights-power-of-public-data-in-uncovering-33b-synthetic-identity-threat/> (дата обращения: 05.08.2026).
- LexisNexis Risk Solutions. Cybercrime Report 2026. Апрель 2026. URL: <https://risk.lexisnexis.com/global/en/about-us/press-room/press-release/20260326-ccr-global-fraud> (дата обращения: 05.08.2026).
- Deloitte Center for Financial Services. Generative AI and Deepfake Banking Fraud Risk. 2024. URL: <https://www.deloitte.com/us/en/insights/industry/financial-services/deepfake-banking-fraud-risk-on-the-rise.html> (дата обращения: 05.08.2026).

Уровень D. Отчёты поставщиков

- Entrust. 2026 Identity Fraud Report. Ноябрь 2025. URL: <https://www.entrust.com/resources/reports/identity-fraud-report> (дата обращения: 05.08.2026).
- Sumsub. Identity Fraud Report 2025–2026. Ноябрь 2025. URL: <https://sumsub.com/fraud-report-2025/> (дата обращения: 05.08.2026).
- iProov. Threat Intelligence Report 2025; исследование восприятия дипфейков (пресс-релиз 12.02.2025). URL: <https://www.iproov.com/press/annual-identity-verification-threat-intelligence-report> ; <https://www.iproov.com/press/study-reveals-deepfake-blindspot-detect-ai-generated-content> (дата обращения: 05.08.2026).
- Smile ID. 2026 Digital Identity Fraud Report. URL: <https://usesmileid.com> (дата обращения: 05.08.2026).

- Group-IB. Weaponized AI: Inside the Criminal Ecosystem. 2026; материалы о голосовых дипфейках. URL: <https://www.group-ib.com/blog/ai-cybercrime-usecases/> (дата обращения: 05.08.2026).
- Shufti. Deepfake / Identity Fraud Index Report 2026 (прогнозные значения помечены поставщиком как 2026E). URL: <https://shuftipro.com/resources/whitepapers-reports/deepfake-identity-fraud-index-report-2026/> (дата обращения: 05.08.2026).
- deepidv. Fraudulent ID & Deepfake Benchmark Report 2026 (иллюстративные данные заинтересованной стороны). URL: <https://www.deepidv.com/fraudulent-identification-benchmark-report-2026> (дата обращения: 05.08.2026).
- Liminal; Unico. Identity Fraud Intelligence: Trends & Insights. Summer 2026 (исходный анализируемый материал; распространяется с ограничением доступа) (дата обращения: 05.08.2026).

Приложение А. Основные термины и метрики

- **Кража идентичности (identity theft)** - завладение персональными данными; **мошенничество с идентичностью (identity fraud)** - их противоправное использование.
- **Синтетическая личность** - составная или полностью искусственная идентичность без реального прототипа, способного заявить о краже.
- **Презентационная атака (PAD)** - предъявление подделки физическому сенсору (фото, экран, маска); **инъекционная атака** - подача манипулированного потока в обход сенсора (виртуальная камера, подмена конвейера).
- **Морфинг** - смешение изображений двух лиц в одно, проходящее сверку с обоими.
- **Захват учётной записи (account takeover)** - получение контроля над действующей учётной записью; **подмена SIM-карты (SIM swap)** - перехват номера для обхода подтверждений.
- **APP-мошенничество** - платёж, санкционированный самой жертвой под воздействием злоумышленника.
- **APCER / BPCER (ISO/IEC 30107-3)** - доля ошибочно принятых атак / доля ошибочно отклонённых добросовестных предъявлений; **FAR / FRR** - вероятности ложного приёма и ложного отказа; **FTXR** - доля неизвлечения (селфи или документа); **FNMR / FMR** - доли ложного несовпадения и ложного совпадения. Корректное сравнение систем возможно только в заявленной операционной точке.
- **Проверяемые учётные данные (verifiable credentials)** - криптографически подписанные утверждения о субъекте с возможностью выборочного раскрытия и отзыва.
- **Доверенный захват (trusted capture)** - получение биометрического образца контролируемым каналом с подтверждённой целостностью, исключающее предоставление заявителем материала с неизвестной историей.

Приложение Б. Официальные таблицы результатов RIVTD (MdTF)

Источник: Maryland Test Facility, официальная страница результатов RIVTD (<https://mdtf.org/rivtd/Results2023>; дата обращения: 05.08.2026). Участники анонимизированы издателем. Обозначения: FTXR - доля неизвлечения; FNMR - доля ложного несовпадения при пороге поставщика с целевым FMR 1:10 000; классификация порога: К - консервативный, П - пермиссивный, О - ожидаемый; NA - не оценено из-за технических сбоев. По треку 3 значения BPCER и APCER - максимум (худший случай) по смартфонам и видам атак.

Б.1. Трек 2: сопоставление селфи с документом (16 систем; 1 633 добровольца)

MTDS	FTXR селфи	FTXR документа	FNMR	Порог
1	0%	<100%	NA	П
2	<1%	<80%	NA	П
3	<1%	<3%	<1%	К
4	<1%	<1%	<1%	К
5	0%	<1%	0%	К
6	0%	<1%	<1%	О
7	<2%	0%	NA	NA
8	<9%	<13%	<100%	О
9	0%	0%	<95%	П
10	<1%	<17%	0%	К
11	0%	<6%	<1%	К
12	<1%	<2%	<9%	О
13	<1%	<100%	NA	NA
14	0%	<13%	<1%	О
15	0%	0%	<1%	П
16	<1%	<2%	<1%	О

Б.2. Трек 3: активные подсистемы детекции презентационных атак (6 систем; 660 участников)

PAD-A	BPCER (макс.)	APCER (макс.)	Удовлетворённость (мин.)	Время транзакции (макс., среднее)
1	12,9%	0,0%	91%	28 с
2	58,6%	6,7%	77%	40 с
3	17,0%	0,0%	78%	33 с
4	11,4%	36,7%	87%	27 с
5	41,5%	31,7%	70%	34 с
6	6,1%	23,3%	89%	32 с

Б.3. Трек 3: пассивные подсистемы детекции презентационных атак (15 систем)

PAD-P	BPCER (макс.)	APCER (макс.)	Время обработки (макс., среднее)
1	16,0%	0,0%	3 с
2	38,0%	21,7%	1 с
3	0,5%	35,0%	4 с
4	0,5%	96,7%	<1 с
5	0,3%	46,7%	<1 с
6	14,7%	3,3%	<1 с
7	0,0%	18,3%	2 с
8	0,3%	3,3%	<1 с
9	8,5%	0,0%	1 с
10	0,0%	98,3%	2 с
11	2,1%	96,7%	1 с
12	0,2%	100,0%	22 с
13	0,2%	38,3%	19 с
14	14,5%	5,0%	28 с
15	31,9%	21,7%	10 с

Приложение В. Юридическая верификация нормативного раздела

Настоящее приложение фиксирует результат сверки юридически значимых утверждений раздела 3.5 с первичными и консолидированными текстами актов. Приложение и раздел 3.5 представляют аналитический обзор и не являются индивидуальным юридическим заключением; для принятия решений с правовыми последствиями необходима консультация квалифицированного юриста в соответствующей юрисдикции.

В.1. Верификационная таблица юридически значимых утверждений

Утверждение раздела 3.5	Первичная норма (точное содержание)	Реквизиты и редакция	Статус	Оценка / редакция
Обязательства для высокорисковых систем приложения III применяются со 02.12.2027	Статья 113 Регламента (ЕС) 2024/1689 в редакции регламента «Цифровой омнибус по ИИ»: обязательства разделов 1–3 главы III для систем по ст. 6(2)/приложению III - с 02.12.2027; по ст. 6(1)/приложению I - с 02.08.2028	Регламент (ЕС) 2024/1689; изменяющий регламент: одобрен ЕП 16.06.2026, Советом 29.06.2026	Действует с 27.07.2026	Подтверждена
Запрещённые практики - с 02.02.2025; модели общего назначения - с 02.08.2025	Статья 113 Регламента (ЕС) 2024/1689 (исходная редакция; Омнибусом не изменялась в этой части)	Регламент (ЕС) 2024/1689, ст. 5, гл. V	Действует	Подтверждена
Биометрическая верификация 1:1 исключена из высокорисковой категории	Приложение III, п. 1(a): не охватывает системы, предназначенные для биометрической верификации, единственная цель которых - подтверждение того, что конкретное физическое лицо является тем, за кого себя выдаёт	Регламент (ЕС) 2024/1689, приложение III, п. 1(a); определения - ст. 3	Действует; применение п. 1(a) как основания высокого риска - с 02.12.2027	Подтверждена
Изъятие для выявления финансового мошенничества ограничено контекстом кредитоспособности	Приложение III, п. 5(b): оценка кредитоспособности и кредитный скоринг физических лиц, за исключением систем, используемых для выявления финансового мошенничества	Регламент (ЕС) 2024/1689, приложение III, п. 5(b)	Действует; применение - с 02.12.2027	Подтверждена

Утверждение раздела 3.5	Первичная норма (точное содержание)	Реквизиты и редакция	Статус	Оценка / редакция
Оспоримость изъятия при неблагоприятных решениях	Прямой нормы нет; вопрос отнесён к толкованию назначения системы (ст. 6(3) и подготавливаемые руководства Комиссии)	Регламент (ЕС) 2024/1689, ст. 6(3)	Толкование	Помечена как аналитическая интерпретация авторов
Кошельки цифровой идентичности - к концу 2026 года	Обязанность предоставить кошельки в течение 24 месяцев после принятия имплементационных актов; акты приняты в конце 2024 года	Регламент (ЕС) 2024/1183 (ст. 5а); имплементационные акты Комиссии, ноябрь–декабрь 2024	Действует; срок исполнения - конец 2026	Подтверждена
Закрытый перечень биометрии ЕБС: изображение лица и запись голоса	Закон определяет виды биометрических персональных данных, используемых для идентификации и аутентификации: изображение лица, полученное с использованием фотовидеоустройств, и запись голоса	Федеральный закон от 29.12.2022 № 572-ФЗ	Действует	Подтверждена
Организации работают с векторами; исходные образцы - в ГИС ЕБС	Размещение образцов в ЕБС; предоставление аккредитованным организациям векторов; запрет хранения исходных изображений и записей вне установленных случаев; особый режим векторов	Федеральный закон № 572-ФЗ, в т. ч. ч. 5 ст. 3 (неприменение к векторам ст. 11 и мер по ст. 19 152-ФЗ)	Действует	Подтверждена; формулировка о «хранении только векторов» заменена на разграничение образцов, векторов и результата проверки
Категории «упрощённая / стандартная / подтверждённая»	Наименования категорий регистрации закреплены на подзаконном и операторском уровне; закон № 572-ФЗ наименования не вводит	Подзаконные акты и материалы оператора ГИС ЕБС	Действует (подзаконный уровень)	Уточнена: добавлена оговорка об уровне закрепления наименований
Онлайн-микрозаймы - идентификация через ЕСИА и ЕБС	Пункт 5.8-2 ст. 7 115-ФЗ: идентификация при дистанционном приёме на обслуживание и аутентификация при дистанционном заключении каждого договора займа с использованием ЕСИА и ЕБС	Федеральный закон от 01.04.2025 № 41-ФЗ; 115-ФЗ	Действует: МФК - с 01.03.2026; МКК - с 01.03.2027; надзорный мораторий для МФК до 01.01.2027 (информация Банка России от 05.05.2026)	Уточнена: добавлены реквизиты, сроки и мораторий

Утверждение раздела 3.5	Первичная норма (точное содержание)	Реквизиты и редакция	Статус	Оценка / редакция
Инициативы о дипфейках и маркировке ИИ-контента	Действующих норм нет; предложения о признании дипфейков отягчающим обстоятельством и о маркировке внесены в Государственную Думу	Законопроекты (июль 2026)	Внесены, не приняты	Уточнена: статус обозначен явно

В.2. Матрица применимости регулирования по типам систем

Тип системы	Регламент ЕС об ИИ	GDPR / 152-ФЗ	eIDAS 2.0	DORA / отраслевые
Биометрическая верификация 1:1 (подтверждение заявленной личности)	Вне высокорисковой категории по п. 1(а) прил. III; общие обязанности (прозрачность, ст. 50) сохраняются	Биометрические данные - специальная категория; требуется правовое основание и согласие	Косвенно (кошелёк как источник атрибутов)	Для финансового сектора - операционная устойчивость (DORA), 115-ФЗ/572-ФЗ
Удалённая биометрическая идентификация (1:N)	Высокорисковая (п. 1(а) прил. III); обязательства с 02.12.2027; отдельные запреты ст. 5	Специальная категория; оценка воздействия обязательна	Косвенно	Отраслевые требования безопасности
Скоринг кредитоспособности с антифрод-компонентом	Высокорисковая по п. 5(b) прил. III; изъятие - только для выявления мошенничества; классификация зависит от назначения (интерпретация - В.1)	Автоматизированные решения - ст. 22 GDPR; права субъекта	Не применимо напрямую	Банковское регулирование; в РФ - требования Банка России
Оркестрация противодействия мошенничеству (сигналы, графы, обмен)	Классификация по назначению и выходным решениям; изъятие п. 5(b) не универсально	Минимизация, псевдонимизация, основания обмена; в РФ - 152-ФЗ	Не применимо напрямую	DORA (ЕС); 115-ФЗ, положения Банка России (РФ)
Идентификация через государственную биометрическую инфраструктуру	Национальная компетенция (вне ЕС); для ЕС - сочетание с eIDAS 2.0	В РФ - 572-ФЗ как специальный закон к 152-ФЗ	В ЕС - кошельки к концу 2026	В РФ - 572-ФЗ, 115-ФЗ (п. 5.8-2 ст. 7), акты ФСТЭК и ФСБ

В.3. Оговорка

Раздел 3.5 и настоящее приложение отражают состояние регулирования на дату обращения к источникам (5 августа 2026 года). Регулирование в рассматриваемой области меняется быстро; перед использованием выводов в договорной, закупочной или комплаенс-работе

необходимо проверить действующие редакции актов и получить консультацию юриста соответствующей юрисдикции.