
ФИНАНСЫ

А.В. МИЛЕНКОВ

доктор экономических наук, доцент,
профессор кафедры мировых финансовых рынков
и финтеха Российского экономического университета имени Г.В. Плеханова

С.В. ФРУМИНА

кандидат экономических наук, доцент,
заведующая кафедрой мировых финансовых рынков
и финтеха Российского экономического университета
имени Г.В. Плеханова; доцент кафедры общественных финансов
Финансового университета при Правительстве Российской Федерации

ФИНАНСОВОЕ КОНСУЛЬТИРОВАНИЕ НА ОСНОВЕ ИИ: ОТ ПРЕОДОЛЕНИЯ ИНФОРМАЦИОННОЙ АСИММЕТРИИ К НОВЫМ РИСКАМ И РЕГУЛЯТОРНЫМ ВЫЗОВАМ

В статье анализируется инициатива Банка России по созданию специализированного ИИ-агента для подбора финансовых продуктов, анонсированная в апреле 2026 г. Исследуются архитектурные и регуляторные подходы к внедрению агентных систем в финансовое консультирование. На основе синтеза зарубежных и российских исследований рассматривается экономическая природа ИИ-посредничества как механизма преодоления информационной асимметрии, который одновременно порождает новые вызовы: технические уязвимости алгоритмов, институциональные риски, манипуляцию интересами пользователей, латентность рынка и проблему распределения ответственности. Особое внимание уделяется эволюции регуляторных подходов за рубежом, выявлению факторов, определяющих успешность внедрения агентных систем, а также анализу мировых прецедентов. Предложена классификация уровней зрелости ИИ-финансовых посредников. Сформулированы рекомендации по управлению ИИ-агентом, направленные на обеспечение баланса между технологической автономией, фидуциарной ответственностью и конкурентной нейтральностью.

Ключевые слова: искусственный интеллект, ИИ-агент, финансовое консультирование, информационная асимметрия, Банк России, технические риски ИИ, поведенческие риски, латентность рынка, ответственность за рекомендации ИИ.

УДК: 336.7, 004.8

EDN: NICHNS

DOI: 10.52180/2073-6487_2026_4_70_95

Введение

В апреле 2026 г. на форуме «ТОЛК», организованном одним из крупнейших российских банков, председатель Центрального банка Российской Федерации Э. Набиуллина анонсировала создание специализированного финансового консультанта (агента) на основе искусственного интеллекта для подбора финансовых продуктов (далее – ИИ-агент). Ключевым условием данного проекта ЦБ РФ назвал полную независимость ИИ-агента от продавцов финансовых продуктов¹. Данная инициатива отражает глобальный тренд трансформации финансового посредничества под воздействием технологий искусственного интеллекта (ИИ), который все чаще рассматривается как инструмент преодоления информационной асимметрии и снижения транзакционных издержек.

Согласно оценкам CFA Institute, архитектура для автономных действий уже существует: BlackRock's Aladdin управляет более \$21 трлн долл. США активов, а Man Group развернула AlphaGPT – агентную ИИ-систему для генерации торговых стратегий². Однако технологическая сложность не гарантирует повышения благосостояния потребителей. В агентных системах критическое значение приобретает не сделка, а дизайн системы – ограничители, пороговые значения и правила эскалации. Соответственно, *вопрос ответственности смещается с уровня исполнения на уровень архитектуры*.

Внедрение ИИ-агентов в финансовое консультирование, несмотря на очевидные преимущества в скорости и доступности, порождает целый ряд принципиально новых вызовов. Среди них – технические риски, такие как «галлюцинации» больших лингвистических моделей (LLM), атаки с отравлением данных (data poisoning attack)³ и неконтролируемый дрейф концепции (concept drift)⁴, которые при этом дополняются и институциональными проблемами. Последние свя-

¹ T-Investments. ТОЛК 2026: Границы между финансовым и нефинансовым секторами размываются. Т-Банк. Инвестиции. 1 апреля 2026 г. <https://www.tbank.ru/invest/social/profile/T-Investments/tolk-2026-granicy-mezhdu-finansovym-i-nefinansovym-sektorom-razmyvayutsya/?author=profile> (дата обращения: 02.04.2026).

² CFA Institute. When AI Trades, Who Is Responsible? March 22 2026 г. <https://rpc.cfainstitute.org/blogs/enterprising-investor/2026/when-ai-trades-who-is-responsible> (дата обращения: 10. 04. 2026).

³ Намеренное внедрение искаженных или вредоносных образцов в обучающую выборку, чтобы модель усвоила неверные закономерности.

⁴ Изменение зависимости между характеристиками объекта (входными признаками) и правильным ответом, из-за которого ранее обученная и не доработанная с учетом этого фактора модель со временем начинает систематически ошибаться, даже если распределение входных данных не изменилось.

заны с ловушкой принципала-агента [15], возникающей из-за конфликта интересов между банками-продавцами (принципалами), разрабатывающими ИИ-агента, и потребителями, чьи интересы он призван защищать. Кроме того, существует риск трансформации рыночной власти (латентность рынка), при котором контроль над дизайном алгоритма ранжирования может стать инструментом нетарифного регулирования. Указанные угрозы усугубляются отсутствием правовой определенности в вопросах распределения ответственности за возможный ущерб, причиненный рекомендациями ИИ.

Несмотря на растущее число исследований, посвященных отдельным аспектам применения ИИ в финансах (кибербезопасность, манипуляция поведением, регуляторный арбитраж), системный анализ институционального дизайна, балансирующего между автономией алгоритма, защитой прав потребителей и конкурентной нейтральностью, остается не завершенным [5]. Особенно это касается национальных инициатив, подобных проекту Банка России, реализующихся в уникальной российской институциональной среде.

Цель настоящей статьи – с помощью системного анализа инициативы по созданию ИИ-агента для подбора финансовых продуктов, рассматриваемой в контексте экономической теории финансового посредничества, выявить риски, определить этапы развития и разработать принципы управления и распределения ответственности за рекомендации, предлагаемые ИИ.

Обзор литературы

Академический дискурс о внедрении искусственного интеллекта в финансовое консультирование формируется на стыке нескольких научных направлений: экономической теории финансового посредничества, институциональной экономики, исследований в области поведенческих финансов и институционального регулирования.

Классические работы, объясняющие природу финансовых посредников как институционального ответа на проблемы информационной асимметрии и трансакционных издержек, развили базовые теоретические положения, способствующие пониманию платформенной экономики [12; 18]. При этом некоторые современные исследования [13] спроецировали существующие положения применительно к ИИ-посредничеству, предлагая классификацию уровней зрелости ИИ-финансовых посредников. В целом функциональное содержание финансового посредничества, по мнению некоторых авторов [3], существенно расширилось, а современные теории рассматривают посредников как системообразующие элементы, обеспечивающие стабильность и развитие финансовой системы.

Работы по цифровой трансформации финансового сектора фокусируются на анализе угроз финансовой стабильности, возникающих при появлении новых институтов (криптовалюты, цифровые валюты), и на необходимости пересмотра макропруденциального регулирования [1; 9; 11]. Вклад в осмысление процесса цифровой трансформации финансовой сферы также вносят труды, содержащие комплексный анализ фундаментальных изменений в денежно-финансовой сфере, вызванных цифровизацией [4; 6; 10].

Вопросы правового регулирования ИИ-консультантов активно разрабатываются в зарубежной литературе. В частности, можно отметить анализ эволюции подхода Комиссии по ценным бумагам и биржам США (SEC – Securities and Exchange Commission) к регулированию прогнозной аналитики данных, демонстрирующий переход от попыток категорического запрета конфликтов интересов к системе, основанной на прозрачности, рыночной дисциплине и пресечении мошенничества, проведенный в 2025 г. [20]. В центре дискуссии находится проблема распространения фидуциарной обязанности на алгоритмические системы. Зарубежные исследователи отмечают, что существующие правила, разработанные исходя из того, что консультантом является человек, не учитывают фундаментальное неравенство сил между инвесторами и платформами, а также новые формы конфликтов интересов, вызванные ИИ-сервисами [16].

Новым направлением, формирующемся в академической литературе, является исследование поведенческих аспектов взаимодействия человека с ИИ-советниками. В работах, посвященных феномену «геймификации» [17], анализируются риски эксплуатации когнитивных уязвимостей пользователей через игровые механики интерфейсов. Эмпирические исследования регуляторов⁵ подтверждают, что генеративные ИИ-чаты могут давать противоречивые, неполные или откровенно неверные финансовые советы, что создает угрозу для потребителей. В российской литературе данный комплекс рисков систематизирован в работе Д.А. Кочергина, который выделяет риски кибербезопасности, дискриминации, «отравления данных» и «галлюцинаций» LLM [6].

Наименее разработанной областью является проблема распределения ответственности за ошибки ИИ-агентов. Международные прецеденты демонстрируют, что регуляторы пока возлагают ответственность на компании – разработчики алгоритмов, но эти случаи касаются в основ-

⁵ FCA (Financial Conduct Authority). Money talks: Lessons from 2 LLM pilots on consumer guidance (Research Note). 28 May 2025. <https://www.fca.org.uk/publications/research-notes/money-talks-lessons-2-llm-pilots-consumer-guidance> (дата обращения: 12.04.2026).

ном вводящих в заблуждение заявлений (AI-washing)^{6,7}. Вопрос о том, как распределить ответственность между оператором платформы, разработчиком алгоритма и поставщиками данных в случае реального ущерба от работы автономного ИИ-агента, остается открытым.

Настоящая работа дополняет существующие исследования по нескольким направлениям. Во-первых, она адаптирует зарубежные концепции уровней зрелости ИИ-финансовых посредников к российскому институциональному контексту. Во-вторых, систематизирует риски внедрения ИИ-агента. В-третьих, предлагает конкретные принципы управления ИИ-агентом, включая распределение ответственности.

Теоретические основы

Экономическая природа финансового консультирования. Как отмечают зарубежные исследователи, на финансовых рынках принятие решений часто ограничено информационной асимметрией, когда одна сторона обладает превосходным знанием о продуктах, рисках или ценах, что создает проблемы неблагоприятного отбора и морального риска [13]. В российской науке теория финансового посредничества также получила развитие в работах, систематизирующих эволюцию научных взглядов на этот институт. В частности, некоторые авторы обобщают три классических подхода к объяснению природы финансовых посредников: трансформационный (преобразование активов по срокам, рискам и ликвидности), транзакционный (снижение издержек взаимодействия между кредиторами и заемщиками) и информационный (преодоление асимметрии информации) [3]. Отмечается, что в условиях цифровизации значение информационного подхода возрастает, а появление новых технологических решений трансформирует традиционные формы посредничества [2], что согласуется с современными представлениями о развитии ИИ-агентов в финансовой сфере.

Традиционные финансовые консультанты выполняли функцию «переводчиков» сложности, снижая издержки поиска, мониторинга и соблюдения регуляторных требований. Однако традиционная

⁶ Retail & Consumer Products Law Observer. (2026). FTC Updates (March 2-6, 2026). Columbia Law School. 9 March 2026. <https://www.retailconsumerproductslaw.com/2026/03/ftc-updates-march-2-6-2026/> (дата обращения: 22.04.2026).

⁷ SEC (U.S. Securities and Exchange Commission). In the Matter of Rimar Capital USA, Inc., Rimar Capital, LLC, Itai Royi Liptz, and Clifford Todd Boro. Admin. Proc. File No. 3-22236. 10 October, 2024. <https://www.sec.gov/enforcement-litigation/distributions-harmed-investors/rimar-capital-usa-inc-rimar-capital-llc-itai-royi-liptz-clifford-todd-boro> (дата обращения: 29.04.2026).

модель консультирования сама порождает системные уязвимости. Выделяются три ключевых недостатка человеческого консультирования: рассогласование стимулов через комиссионно-ориентированные продажи, субъективные искажения в суждениях и ограниченную масштабируемость, которая сужает возможности для доступа к профессиональным консультациям [13]. Именно эти уязвимости призваны преодолеть алгоритмические системы.

Однако алгоритмические решения, устраняя одни проблемы, порождают новые, специфические для самой технологии. Если ошибка человека-консультанта локальна и индивидуальна, то алгоритм, обученный на данных, содержащих системные искажения, будет транслировать их во всех своих решениях, многократно усиливая эффект. В отличие от традиционного человеческого консультирования здесь возникают риски, связанные с масштабированием ошибок, заложенных на этапе обучения модели. Это порождает качественно новые угрозы, ранее в практике финансового консультирования не встречавшиеся.

Фидуциарная обязанность как нормативный ориентир. Фидуциарная обязанность, закрепленная в американском Законе об инвестиционных консультантах⁸, опирается на два ключевых элемента: обязанность заботы (duty of care) и обязанность лояльности (duty of loyalty). Согласно разъяснению SEC⁹ и анализу экспертов Columbia Law School [20], именно обязанность лояльности накладывает на консультанта требование ставить интересы клиента выше собственных и урегулировать все возможные конфликты интересов, которые могут повлиять на объективность его рекомендаций. Типичный способ такого урегулирования – это полное и прозрачное раскрытие информации, позволяющее клиенту принять осознанное решение. Фидуциарный стандарт не ограничивается моментом рекомендации и сохраняет силу в течение всего срока взаимодействия, что предъявляет особые требования к системам ИИ-консультирования в части непрерывного контроля качества предлагаемых решений.

Агентный ИИ: сдвиг ответственности на архитектуру. Различие между традиционными алгоритмами и современными агентными системами ИИ принципиальное: традиционные алгоритмы и анали-

⁸ U.S. Congress. (1940). Investment Advisers Act of 1940. Public Law 76-768, 15 U.S.C. §§ 80b-1 et seq. <https://www.govinfo.gov/content/pkg/USCODE-2023-title15/pdf/USCODE-2023-title15-chap2D-subchapII.pdf> (дата обращения: 22.04.2026).

⁹ SEC (U.S. Securities and Exchange Commission). (2019). Commission Interpretation Regarding Standard of Conduct for Investment Advisers. Release No. IA-5248. July 12 2019. <https://www.federalregister.gov/documents/2019/07/12/2019-12208/commission-interpretation-regarding-standard-of-conduct-for-investment-advisers> (дата обращения: 02.05.2026).

тика предоставляют информацию для принятия решений, тогда как агентный ИИ исполняет их. Эти системы не просто рекомендуют – они оценивают, принимают решения и действуют в рамках заранее определенных ограничений¹⁰. Этот сдвиг имеет фундаментальные последствия для распределения ответственности. В традиционной модели решение равно сделке: человек инициирует действие, ответственность зафиксирована в точке исполнения. В агентных системах решение равно дизайну системы: система инициирует действие, ответственность перемещается на архитектуру, то есть агентный ИИ не устраняет подотчетность, а перемещает ее.

Инициатива Банка России

ИИ-агент как инструмент преодоления асимметрии. Инициатива Банка России исходит из корректной диагностики проблемы: сложность финансовых продуктов и асимметрия информации между продавцами и потребителями создают системные угрозы. Председатель Банка России Э. Набиуллина на форуме «ТОЛК» отметила, что платформенная экономика и экосистемы будут развиваться дальше и новые технологии способны снизить издержки и повысить прозрачность финансовых услуг. Однако для реализации этого потенциала необходимо обеспечить, чтобы потребители в полной мере получали благо от технологического прогресса¹¹. Это соответствует современным представлениям о потенциале ИИ в финансовом секторе.

Следует отметить, что в отличие от сегмента профессиональных консультационных услуг для бизнеса, автоматизированное финансовое консультирование для населения пока не получило в России широкого распространения. При этом глобальный рынок робоэдвайзинга, по прогнозам Research and Markets, к 2029 г. может вырасти до 470,9 млрд долл. США¹².

Нормативная база Банка России в сфере ИИ. Инициатива по созданию ИИ-агента для подбора финансовых продуктов формируется

¹⁰ CFA Institute. When AI Trades, Who Is Responsible? March 22 2026. <https://rpc.cfainstitute.org/blogs/enterprising-investor/2026/when-ai-trades-who-is-responsible> (дата обращения: 02.05.2026).

¹¹ T-Investments. ТОЛК 2026: Границы между финансовым и нефинансовым секторами размываются. Т-Банк. Инвестиции. 1 апреля 2026 г. <https://www.tbank.ru/invest/social/profile/T-Investments/tolk-2026-granicy-mezhdu-finansovym-i-nefinansovym-sektorom-razmyvayutsya/?author=profile> (дата обращения: 02.04.2026).

¹² Reuters. (2025, September 25). 'ChatGPT, what stocks should I buy?' AI fuels boom in robo-advisory market. CNBC TV18. <https://www.cnbc.tv18.com/technology/chatgpt-what-stocks-should-i-buy-ai-fuels-boom-in-robo-advisory-market-19690479.htm> (дата обращения: 02.05.2026).

в контексте сложившейся нормативной среды, которую Банк России выстраивает с 2025 г. Регулятор последовательно придерживается подхода «мягкого регулирования» в сфере искусственного интеллекта, что находит отражение в принятых документах.

Основой нормативной базы для ИИ-агента можно считать «Кодекс этики в сфере разработки и применения искусственного интеллекта на финансовом рынке»¹³, разработанный Банком России в июле 2025 г. Документ, носящий рекомендательный характер, декларирует принципы человекоцентричности, прозрачности и ответственного управления рисками. Параллельно ЦБ РФ запустил регулятивную «песочницу» для тестирования ИИ-решений в контролируемых условиях, подтверждая, что использование ИИ не снимает с финансовых организаций ответственности за права клиентов.

Показательно, что ЦБ РФ осознает границы использования технологий. Так, в мае 2025 г. Э. Набиуллина в своем выступлении на форуме подчеркнула, что ИИ не заменит человека в стратегических вопросах, особенно в макроэкономическом анализе и прогнозировании. В то же время регулятор видит потенциал ИИ в повышении прозрачности, снижении издержек и улучшении клиентского опыта¹⁴.

Таким образом, к моменту анонсирования ИИ-агента Банк России уже сформировал нормативный фундамент: закреплены этические принципы, определены требования к информированию клиентов, созданы механизмы тестирования и зафиксирована ответственность организаций. Однако эти документы носят рекомендательный характер и не содержат жестких требований к архитектуре независимых ИИ-агентов, что оставляет пространство для дальнейшей их конкретизации в рамках рассматриваемой инициативы.

Международный контекст. В 2023 г. Комиссия по ценным бумагам и биржам США (SEC) вынесла на обсуждение проект правил «Конфликты интересов, связанные с использованием прогнозной аналитики данных брокерами-дилерами и инвестиционными консультантами»¹⁵. Проект предлагал обязать посредников выявлять любые конфликты

¹³ Банк России. Кодекс этики в сфере разработки и применения искусственного интеллекта на финансовом рынке: Информационное письмо от 09.07.2025 № ИН-016-13/91. <https://www.cbr.ru/Crosscut/LawActs/File/10045> (дата обращения: 02.05.2026).

¹⁴ Известия. Набиуллина рассказала о применении ИИ в работе Банка России. 20 мая 2025. <https://iz.ru/1889213/2025-05-20/nabiullina-rasskazala-o-primenenii-ii-v-rabote-banka-rossii> (дата обращения: 02.04.2026).

¹⁵ SEC (U.S. Securities and Exchange Commission). (2023). Conflicts of Interest Associated with the Use of Predictive Data Analytics by Broker-Dealers and Investment Advisers. Proposed Rule Release No. 34-97990. <https://www.federalregister.gov/documents/2023/08/09/2023-16377/conflicts-of-interest-associated-with-the-use-of-predictive-data-analytics-by-broker-dealers-and> (дата обращения: 22.04.2026).

интересов, связанные с использованием «охватываемых технологий» («аналитическая, технологическая или вычислительная функция, алгоритм, модель, корреляционная матрица или аналогичный метод или процесс, который оптимизирует, предсказывает, направляет, прогнозирует или определяет инвестиционное поведение или результаты инвестора»), и затем устранять или нейтрализовывать их, фактически исключая раскрытие информации как способ урегулирования конфликта. Критики проекта SEC указывали на неоправданную широту определения «охватываемых технологий», разрыв с устоявшейся практикой раскрытия информации и заниженную оценку издержек [20], и в июне 2025 г. SEC отозвала проект¹⁶.

Альтернативный подход предлагает гибкую трехкомпонентную систему: прозрачность, рыночную дисциплину и жесткое пресечение мошенничества [20]. Наглядным примером служит «дело Cleo AI». В январе 2026 г. FTC оштрафовала компанию на \$17 млн долл. США за ложные обещания ИИ-бота по экономии средств и за использование игровых механик, которые подталкивали пользователей к накоплениям, но затрудняли вывод денег¹⁷. Это дело демонстрирует, что регуляторы квалифицируют геймификацию как нарушение прав потребителей.

Другой ценный опыт связан с пилотными проектами FCA в Великобритании в 2025 г. с использованием GPT-3.5 и GPT-4¹⁸. Они показали, что LLM хорошо упрощают финансовые концепции для неопытных пользователей, но их эффективность сильно зависит от контекста, а проверка результатов требует сочетания автоматизации и человеческого контроля. FCA также зафиксировала высокий спрос на ИИ-ассистентов.

В том же году SEC привлекла к ответственности компании Delphia и Global Predictions за «AI-вошинг», т. е. за введение клиентов в заблуждение ложными заявлениями об использовании ИИ. Фирмы заявляли, что «активно используют ИИ» и «полагаются на ИИ», в то время как на

¹⁶ SEC (U.S. Securities and Exchange Commission). Conflicts of Interest Associated with the Use of Predictive Data Analytics by Broker-Dealers and Investment Advisers; Withdrawal. Release No. 34-97990; IA-6279. June 12 2025. <https://www.sec.gov/rules-regulations/2025/06/s7-12-23> (дата обращения: 21.04.2026).

¹⁷ Retail & Consumer Products Law Observer. FTC Updates (March 2-6, 2026). Columbia Law School. 9 March 2026. <https://www.retailconsumerproductslaw.com/2026/03/ftc-updates-march-2-6-2026/> (дата обращения: 22.04.2026).

¹⁸ FCA (Financial Conduct Authority). Money talks: Lessons from 2 LLM pilots on consumer guidance (Research Note). 28 May 2025. <https://www.fca.org.uk/publications/research-notes/money-talks-lessons-2-llm-pilots-consumer-guidance> (дата обращения: 22.04.2026).

самом деле не применяли таких технологий¹⁹. Этот случай иллюстрирует растущий интерес к достоверности заявлений об использовании ИИ и важность их регуляторной проверяемости.

Перспективы развития финансового консультирования и дорожная карта

Демократизация, макропруденциальный анализ и технологический суверенитет. Зарубежные исследователи отмечают, что робоэдвайзеры продемонстрировали потенциал алгоритмических инструментов для демократизации доступа к инвестиционным консультациям, повышения эффективности за счет автоматизации и обеспечения дисциплины в ребалансировке портфелей [13]. Отечественный проект, при условии обеспечения независимости, может существенно расширить доступ к финансовому консультированию для населения, особенно для сегментов, которые традиционно не обслуживаются человеческими консультантами из-за низкой маржинальности, что в терминах концепции финансового развития отвечает расширению финансовой доступности [7]. Это позволит снизить социальное неравенство в доступе к качественным финансовым услугам и повысить финансовую грамотность широких слоев населения.

ИИ-агент, агрегирующий информацию о поведении потребителей, может стать источником ценных данных для макропруденциального анализа, позволяя Банку России более точно оценивать угрозы, связанные с поведенческими факторами. Результаты эмпирических исследований подтверждают, что нейронные сети способны выделять из «шума» социальных сетей индикаторы коллективных настроений, коррелирующие с будущей волатильностью торговой активности и прогнозировать ее иррациональные всплески [19].

Разработка собственного ИИ-агента выходит за рамки технологического эксперимента и приобретает стратегическое значение. В условиях санкционных ограничений появление отечественного ИИ-решения обеспечивает снижение зависимости от иностранных сервисов в критической финансовой инфраструктуре. Более того, адаптация алгоритмов к национальной специфике (налогообложение, валютное регулирование) может повысить их релевантность для российских потребителей. Кроме того, создание суверенного ИИ-агента

¹⁹ Herbert Smith Freehills. SEC Takes Action Against 'AI Washing,' Fines Two Investment Advisers for Misrepresenting Artificial Intelligence Use. HSF Kramer Blog. 26 June 2025. <https://www.hsfkramer.com/insights/2024-03/sec-takes-action-against-ai-washing-fines-two-investment-advisers-for-misrepresenting-artificial-intelligence-use> (дата обращения: 29.04.2026).

может стимулировать развитие национальной школы в области финансового ИИ и подготовку профильных кадров. Таким образом, инициатива Банка России представляет собой не только инструмент защиты прав потребителей, но и проект, связанный с укреплением технологического суверенитета России в финансовой сфере.

Дорожная карта: уровни зрелости ИИ-агента. На основе предложенной некоторыми исследователями классификации ИИ-финансовых посредников можно выделить этапы развития отечественной инициативы (см. табл.).

Таблица

Уровни зрелости ИИ-агента

Уро- вень	Название	Характеристика	Возможный целевой горизонт для России
1	Калькулятор	Предоставление информации в унифицированном формате, отсутствие персонализации	Реализован
2	Рекомен- дательная система	Персонализированные варианты на основе профиля клиента	Реализован (коммерческие агрегаторы); для инициативы ЦБ РФ – адаптация под требования независимости
3	Агент с ограни- ченными полномочиями	Автоматические действия в рамках заданных ограничений, разделение полномочий	Частично реализован (отдельные банки); для инициативы ЦБ РФ – адаптация под требования независимости, тиражирование и стандартизация
4	Адаптивный планировщик	Учет изменений внешней среды	Перспективный уровень
5	Интеллекту- альный пла- нировщик	Предвосхищение потребностей клиента	Долгосрочная перспектива

Источник: составлено авторами на основе: [13].

Анализируя таблицу, можно сказать, что уровни 2 и 3 уже в значительной степени реализованы на российском рынке коммерческими агрегаторами и отдельными банками. Уникальность инициативы Банка России заключается не в создании технологии «с нуля», а в обеспечении независимости от продавцов финансовых продуктов, в создании системы сдержек и противовесов и достижении системной значимости (массовое использование ИИ-агента – инициативы Банка Рос-

сии, выступающего в качестве официального канала рекомендаций по финансовому консультированию). По нашему мнению, достижение более высоких уровней зрелости ИИ-агента (4–5) потребует не только повышения уровня развития информационно-коммуникационных технологий, но и решения ряда институциональных проблем, включая когнитивную независимость, устойчивое финансирование и нормативно закреплённую ответственность.

Анализ рисков внедрения ИИ-агента

Технические риски и угрозы внешней зависимости. Внедрение ИИ-агента сопряжено с рядом *технических рисков*, присущих алгоритмическим системам. В широком контексте цифровой трансформации финансового сектора эти уязвимости становятся частью системных вызовов, требующих переосмысления устоявшихся подходов к регулированию и управлению рисками [10]. В качестве главных возможных технических рисков выделим следующие.

Риск дрейфа управления. Прецедент Two Sigma Investments, урегулировавшей претензии SEC на \$90 млн долл. США после того, как исследователь в течение почти четырех лет модифицировал торговые алгоритмы без адекватного контроля, иллюстрирует ключевую угрозу агентных систем – постепенное расширение автономии алгоритма без соответствующей эволюции надзора²⁰. Аналогичные выводы представлены в обзоре Управления по финансовому поведению Великобритании (FCA). В 2025 г. регулятор выявил устойчивые слабости, такие как устаревшие политики, нечеткие структуры подотчетности и недостаточное тестирование в фирмах, использующих алгоритмический трейдинг²¹.

Риски, связанные с качеством и безопасностью данных. Российские эксперты отмечают, что внедрение ИИ в финансовой сфере порождает комплекс технических рисков, связанных с данными и алгоритмами: генеративные модели расширяют возможности злоумышленников по созданию фишинговых писем, вредоносного ПО и захвату устройств, что создает угрозу кражи данных, вымогательства и мошенничества; алгоритмы способны воспроизводить и усиливать искажения из исходных данных, что ведет к дискриминации; возможны атаки

²⁰ CFA Institute. When AI Trades, Who Is Responsible? 22 March, 2026. <https://rpc.cfainstitute.org/blogs/enterprising-investor/2026/when-ai-trades-who-is-responsible> (дата обращения: 22.04.2026).

²¹ FCA (Financial Conduct Authority). Multi-firm review of algorithmic trading controls: high-level observations. 21 August, 2025. <https://www.fca.org.uk/publications/multi-firm-reviews/algorithmic-trading-controls-high-level-observations> (дата обращения: 12.04.2026).

с отравлением данных, когда злоумышленники вмешиваются в обучающие выборки; генеративные модели подвержены «галлюцинациям» – генерации правдоподобных, но неверных ответов [6].

Риск устаревания данных. При централизованной инфраструктуре (например, если агент работает на мощностях ЦБ РФ) возникает угроза задержек в обновлении информации. В высокочастотной рознице (вклады, кредитные карты) условия могут меняться достаточно часто, и централизованный агент рискует предлагать пользователям предложения, которые уже недоступны или изменились, что подрывает доверие к системе.

Использование зарубежных ИИ-решений в финансовой сфере создает дополнительные геополитические риски (*риски внешней зависимости*): несанкционированный доступ к данным граждан, блокировка сервисов из-за санкций и неявные предубеждения в алгоритмах, не соответствующие российским экономическим реалиям.

Институциональные риски и латентность рынка. Ключевая характеристика ИИ-агента, анонсированного Банком России, – его полная независимость от продавцов финансовых продуктов. Однако, как известно, независимость всегда требует устойчивого механизма финансирования, который неизбежно порождает стимулы и ограничения, способствующие реализации специфических институциональных рисков. К ним можно отнести следующие риски.

Риск неопределенности модели финансирования. Отсутствие публично закреплённой модели финансирования усиливает неопределённость и подрывает доверие к проекту.

Риск недоинвестирования. Государственное финансирование чревато риском недоинвестирования, устаревания технологий и потенциальными атаками со стороны участников, теряющих доходы.

Финансирование проекта за счет банков, в свою очередь, порождает классическую ловушку «принципал-агент», когда банки платят за инструмент, который перераспределяет их клиентов к конкурентам. В этой связи можно выделить дополнительные риски.

Риск когнитивной зависимости. Участие банков-продавцов в разработке алгоритма означает, что данные для обучения могут быть неявно смещены в их пользу. Это институциональная проблема, которая возникает каждый раз, когда оценщик (ИИ-агент) создается при участии оцениваемых (банков).

Риск оппортунистического поведения. Банки-участники могут использовать доступ к внутренней логике алгоритма для создания дополнительных конкурентных преимуществ.

Риск саботажа качества данных. Банки будут предоставлять ИИ-агенту упрощенные или «очищенные» версии условий, скрывая реальные подводные камни. При этом формальные требования к рас-

крытию информации могут соблюдаться, но полезность этих данных для качественного анализа ИИ-агентом окажется сниженной.

Таким образом, экономика агента оказывается центральным вопросом устойчивости проекта. Отсутствие ответа на вопрос об источнике финансирования объективности ставит под сомнение любое утверждение о «полной независимости», если оно не подкреплено соответствующей институциональной базой.

Следует отметить, что в исходной формулировке инициативы Банка России ИИ-агент задуман как инструмент защиты прав потребителей и преодоления информационной асимметрии. Однако системный анализ, опирающийся на теорию двухуровневых рынков [12; 18] и на исследования цифровых платформ [14], позволяет выявить долгосрочные *риски латентности рынка* (скрытое свойство рыночной структуры трансформироваться из конкурентной в платформенно-зависимую при достижении ИИ-агентом системной значимости). В отличие от традиционных рынков, где конкуренция происходит преимущественно по цене и качеству, на двухуровневых рынках платформа-посредник может приобрести рыночную власть за счет контроля над взаимодействием между группами пользователей.

Если ИИ-агент станет системным, то есть начнет направлять значимые клиентские потоки, а потребители будут воспринимать его как авторитетный и независимый источник рекомендаций, структура конкуренции на финансовом рынке, по нашему мнению, может претерпеть фундаментальные изменения. Банки могут перестать конкурировать преимущественно условиями и начать конкурировать ранжированием в ИИ-агенте. Как отмечают некоторые исследователи, без четких гарантий цифровые платформы рискуют воспроизвести давние рыночные неэффективности: информационную асимметрию, несогласование стимулов и системную хрупкость [13]. В данном случае риск заключается в том, что ИИ-агент, призванный решить проблему асимметрии, сам становится новой точкой концентрации рыночной власти.

Российские эксперты обращают внимание на близкий по смыслу риск – зависимость от ограниченного числа поставщиков моделей ИИ. Использование многими финансовыми организациями одних и тех же технологических решений и данных приводит к синхронизации принимаемых решений. Различные участники рынка автоматически совершают схожие действия, что создает эффект повышенной волатильности и риска манипуляций [6]. Этот вывод имеет двойственное значение для российской инициативы. С одной стороны, если алгоритмы ИИ-агента будут непрозрачны, он сам станет источником «стадного» поведения банков, когда все начнут действовать одинаково, угадывая алгоритм. С другой стороны, сам факт наличия независимого

государственного агента снижает зависимость банков от иностранных коммерческих ИИ-решений, делая рынок более устойчивым и разнообразным.

Риски латентности рынка имеют два измерения.

Первое измерение – техническое, создающее *риск регуляторного арбитража через алгоритм*. Алгоритм ранжирования ИИ-агента должен опираться на определенные критерии и веса. Если эти критерии не являются абсолютно прозрачными и машиночитаемыми, возникает пространство для «регуляторного арбитража через ИИ»: банки будут оптимизировать свои продукты не под реальные потребности потребителей, а под критерии ранжирования агента. В конечном счете это может привести к ситуации, когда формальные характеристики продукта улучшаются, но реальное качество для потребителя может не расти.

Второе измерение – институциональное, создающее *риск конфликта регуляторных функций*. Как отмечалось выше, в агентных системах критическое решение – это дизайн системы, то есть ограничители, пороговые значения и правила эскалации. Если ИИ-агент становится системным, контроль над его дизайном превращается в инструмент рыночного регулирования, сопоставимый по значимости с денежно-кредитной политикой. Если Банк России окажется в роли оператора алгоритмической системы, которая перераспределяет клиентские потоки, то это создает новые вызовы для самого регулятора. Во-первых, в отличие от коммерческого оператора регулятор не может принимать на себя полную финансовую ответственность за возможные убытки потребителей – это противоречило бы его природе и создавало бы неприемлемые риски для самого ЦБ РФ. Во-вторых, возникает ситуация, когда центральный банк одновременно является и регулятором, устанавливающим правила и осуществляющим надзор, и оператором, управляющим алгоритмом, который эти правила реализует на практике. Это порождает угрозы неопределенности подотчетности (неясно, кто и перед кем отвечает за возможные ошибки) и создает предпосылки для использования алгоритма ранжирования в целях нетарифного регулирования (скрытого перераспределения клиентских потоков в пользу определенных участников рынка). Следовательно, институциональный дизайн ИИ-агента должен исключать саму возможность совмещения регуляторных и операторских функций в одном лице.

Риски поведенческого воздействия и уязвимости данных для макропруденциального анализа. Помимо указанных выше рисков, внедрение ИИ-агента в финансовое консультирование порождает специфическую категорию рисков, связанных с возможностью манипуляции поведением и интересами пользователей (*риски поведенческого*

воздействия). В академической литературе эти риски описываются как «поведенческая эксплуатация» (behavioral exploitation), когда дизайн платформы намеренно использует когнитивные уязвимости пользователей [17], и как «алгоритмическая манипуляция» (algorithmic manipulation) [8; 13].

Некоторые исследователи выделяют три основных канала манипуляции в ИИ-системах финансового консультирования [16], формирующих соответствующие риски.

Риск геймификации интерфейсов. Платформы используют игровые механики (уведомления, анимации, визуальные эффекты, индикаторы успеха) для стимулирования более активного и рискованного поведения пользователей. Зависимость Robinhood от геймифицированных интерфейсов и платы за поток заказов иллюстрирует опасности поведенческой манипуляции: в ходе событий вокруг акций GameStop в 2021 г. тактика вовлечения платформы активно направляла внимание пользователей на волатильные акции, а торговля опционами резко выросла среди неопытных инвесторов [13].

Риск манипуляции через непрозрачное ранжирование. ИИ-агент, даже формально независимый, может использовать методы «подталкивания» пользователя в направлении к определенным продуктам. Если критерии ранжирования не являются абсолютно прозрачными и не утверждены регулятором, возникает риск того, что агент будет систематически предлагать продукты, которые приносят выгоду его разработчикам или партнерам, а не пользователю. В конечном счете это может превратить инструмент защиты потребителя в инструмент «алгоритмического внушения», где пользователь субъективно сохраняет свободу выбора, но фактически следует предопределенной траектории [16].

Риск манипуляции через гиперперсонализацию. ИИ-системы, анализирующие поведение пользователя, могут выявлять его эмоциональные уязвимости (тревожность, склонность к импульсивным решениям, финансовую неграмотность) и использовать их для продвижения продуктов. Отмечается, что клиенты, вовлеченные в неэффективные или неточные взаимодействия с ИИ-агентами, могут терять время, принимать неверные инвестиционные решения, а платформа может извлекать выгоду из эмоционального расстройства инвесторов [16]. В контексте российской инициативы это означает, что даже независимый ИИ-агент, если он не имеет фидуциарной обязанности перед пользователем, может использовать гиперперсонализацию не для защиты интересов клиента, а для максимизации его вовлеченности.

Использование ИИ-агента в качестве источника данных для макропруденциального анализа, рассматриваемое нами выше как возможная перспектива, сопряжено с рядом рисков (*риски использования дан-*

ных для макропруденциального анализа). Как отмечают некоторые эксперты, цифровая трансформация финансового сектора порождает «турбулентность угроз финансовой стабильности», связанную с появлением новых институтов и инструментов, а скорость распространения таких явлений способна дестабилизировать финансовую систему, что требует пересмотра традиционных подходов к оценке системных угроз, не учитываемых в существующих моделях [1].

Риск нерепрезентативности данных. Пользователи ИИ-агента могут составлять специфическую выборку, не отражающую всю совокупность потребителей финансовых услуг. Например, на начальных этапах это могут быть более молодые, финансово грамотные и технологически продвинутые пользователи. Экстраполяция их поведения на всех потребителей может привести к ошибочным выводам о системных рисках.

Риск манипуляции данными со стороны банков. Как отмечалось ранее, у банков есть стимулы к сокрытию ключевых характеристик продуктов. Если ИИ-агент становится источником данных для регулятора, банки могут сознательно искажать информацию, предоставляемую агенту, чтобы повлиять на оценки макропруденциальных рисков в свою пользу.

Риск неправильной интерпретации данных. Сложность алгоритмов ИИ-агента и агрегированных данных может создавать иллюзию точности и объективности. Регулятор, не обладая полным пониманием внутренней логики агента (проблема «черного ящика»), может сделать неверные выводы о корреляциях и причинно-следственных связях, что приведет к принятию ошибочных макропруденциальных мер.

Итак, проведенный анализ рисков внедрения ИИ-агента позволяет выделить шесть ключевых групп: технические, внешние, институциональные, поведенческие, а также риски латентности рынка и использования данных для макропруденциального анализа (см. рис.). Многообразие выявленных угроз обуславливает необходимость комбинированного подхода, объединяющего архитектурные решения, регуляторные нормы и институциональные ограничения. Системный учет всех категорий рисков на этапе проектирования ИИ-агента является необходимым условием его устойчивости и сохранения доверия со стороны потребителей и участников финансового рынка.

Представленная классификация рисков носит качественный характер. Количественная оценка вероятности реализации каждого риска и потенциального экономического ущерба в настоящем исследовании не проводится. По мере развертывания пилотной версии и накопления статистических данных о поведении пользователей и банков-контрагентов целесообразна разработка соответствующих имитационных и эконометрических моделей.



Источник: составлено авторами.

Рис. Риски внедрения ИИ-агента

Ответственность за ошибки ИИ-агента

По нашему мнению, одним из наиболее сложных и недостаточно проработанных аспектов инициативы Банка России является вопрос о том, кто будет нести реальную ответственность за ошибки ИИ-агента и фактические убытки потребителей, воспользовавшихся его рекомендациями. В действующей нормативной базе (включая Кодекс этики ЦБ РФ) этот вопрос не урегулирован.

В Стратегии развития искусственного интеллекта до 2030 г. (в редакции Указа Президента РФ от 15.02.2024 г. № 124) закреплён принцип – ответственность за результаты работы ИИ не может быть делегирована самому искусственному интеллекту; она сохраняется за человеком или организацией, использующей ИИ. Однако этот принцип не конкретизирует, какое именно лицо (разработчик алгоритма, оператор системы, банк-партнер или сам Банк России) отвечает в случае причинения убытков.

Как отмечалось выше, алгоритмические системы могут принимать решения, которые невозможно объяснить на доступном для человека уровне. Это создает фундаментальную проблему для системы ответственности: если регулятор, суд или сам потребитель не могут понять, почему ИИ-агент дал ту или иную рекомендацию, оспорить ее или доказать наличие ошибки становится крайне сложно. Даже при формальном закреплении ответственности за оператором, практическая реализация права на обжалование может оказаться затруднительной из-за невозможности ретроспективного анализа логики работы алгоритма. Возможным решением может быть внедрение методов «объяснимого» ИИ на этапе архитектурного проектирования агента.

В мировой практике можно выделить следующие прецеденты, связанные с определением ответственных лиц. В деле *Rimar Capital 2024 г.* SEC обвинила компанию и ее руководителей в привлечении около \$4 млн долл. США от 45 инвесторов на основе ложных заявлений об использовании ИИ для автоматизированного трейдинга, в то время как такой технологии не существовало²². В деле *Cleo AI 2026 г.* Федеральная торговая комиссия США (FTC) привлекла к ответственности компанию-разработчика за введение потребителей в заблуждение относительно возможностей ИИ-бота по экономии денег²³. В Великобритании после испытаний ИИ-чатов регуляторы (FCA) в отчете «Money Talks» указали на необходимость четкого распределения ответственности при использовании сторонних ИИ-инструментов в финансовых услугах²⁴.

²² SEC (U.S. Securities and Exchange Commission). In the Matter of Rimar Capital USA, Inc., Rimar Capital, LLC, Itai Royi Liptz, and Clifford Todd Boro. Admin. Proc. File No. 3-22236. 10 October, 2024. <https://www.sec.gov/enforcement-litigation/distributions-harmed-investors/rimar-capital-usa-inc-rimar-capital-llc-itai-royi-liptz-clifford-todd-boro> (дата обращения: 02.05.2026).

²³ Retail & Consumer Products Law Observer. FTC Updates (March 2-6, 2026). Columbia Law School. 9 March 2026. <https://www.retailconsumerproductslaw.com/2026/03/ftc-updates-march-2-6-2026/> (дата обращения: 02.05.2026).

²⁴ FCA (Financial Conduct Authority). Money talks: Lessons from 2 LLM pilots on consumer guidance (Research Note). 28 May, 2025. <https://www.fca.org.uk/publications/research-notes/money-talks-lessons-2-llm-pilots-consumer-guidance> (дата обращения: 02.05.2026).

На наш взгляд, институциональная модель, основанная на разделении регуляторных и операторских функций, должна формировать цепочку ответственности, включающую как минимум четыре потенциальных субъекта:

- Банк России – инициатор создания ИИ-агента, регулятор, осуществляющий надзор за деятельностью оператора платформы и утверждающий нормативные требования к работе ИИ-агента;
- оператор платформы – отдельная юридическая структура (государственная IT-компания или уполномоченная НКО), непосредственно управляющая алгоритмом ранжирования, обеспечивающая его работу и подотчетная регулятору;
- банки – участники разработки – организации, предоставляющие данные и технологическую экспертизу на этапе создания ИИ-агента (без доступа к актуальной версии алгоритма и возможности влиять на ранжирование в реальном времени);
- иные банки-партнеры – поставщики данных о своих продуктах, чьи предложения ранжируются ИИ-агентом.

Для обеспечения дополнительного внешнего мониторинга деятельности оператора может быть создан общественный совет. Его функции, в отличие от регулятора, могут носить не властный, а контрольно-рекомендательный характер: регулярный аудит конкурентной нейтральности алгоритма, оценка жалоб потребителей, подготовка ежегодного публичного отчета. В отличие от внутренних механизмов оператора (разделение разработки, валидации и эксплуатации) общественный совет представляет интересы общества, потребителей и добросовестных банков, обеспечивая публичную прозрачность и доверие к системе.

В отсутствие четкого нормативного регулирования возможны две крайности. Во-первых, «размывание ответственности», когда каждый участник (оператор, банки – участники разработки, банки-партнеры) ссылается на других, и потребитель не может получить компенсацию. Во-вторых, возложение всей ответственности на Банк России, как на регулятор и инициатор создания агента (без выделения отдельного оператора), что создает для него неприемлемые финансовые и репутационные риски. Рассматриваемое в данной статье разделение функций (ЦБ – регулятор и надзор, оператор – отдельная структура) позволяет избежать обеих крайностей. В этой связи, на наш взгляд, целесообразно закрепить нормативно следующие принципы распределения ответственности:

1) оператор ИИ-агента (отдельная юридическая структура, не совмещающая надзорные функции с функциями управления алгоритмом) должен быть подотчетен Банку России как регулятору в части корректной работы алгоритма ранжирования и достоверности предоставляемой им информации. Это означает, что оператор обязан

организовать систему контроля и аудита, обеспечить прозрачность критериев ранжирования, публично отчитываться о работе агента и принимать меры по устранению выявленных нарушений. Банк России при этом осуществляет надзор за деятельностью оператора, но не участвует в оперативном управлении алгоритмом;

2) финансовую ответственность за убытки потребителей, причиненные рекомендациями ИИ-агента, должны нести:

- банки-партнеры – за убытки, возникшие из-за недостоверной или неполной информации о своих продуктах;
- банки – участники разработки – за убытки, вызванные дефектами, внесенными на этапе создания алгоритма (ошибки в архитектуре, некорректный выбор обучающих данных, дефекты валидации), при условии, что будет доказана их вина в форме умысла или грубой неосторожности;
- специальный компенсационный фонд (формируемый за счет обязательных отчислений банков-партнеров) – за убытки, вызванные непредсказуемым поведением алгоритма на этапе эксплуатации, когда невозможно установить вину конкретного участника (например, нештатная реакция обученной модели на изменение внешних условий);
- оператор должен нести субсидиарную ответственность перед компенсационным фондом: если вина оператора в причинении убытков будет доказана, фонд имеет право регрессного требования к оператору на сумму выплаченной компенсации.

3) потребитель должен иметь право на обжалование рекомендации и получение компенсации убытков из указанных источников.

Заключение и рекомендации

Ключевые выводы исследования.

1. В отличие от традиционных алгоритмов и аналитики, которые только формируют решения, агентный ИИ самостоятельно исполняет их. В результате ответственность смещается с точки совершения сделки на уровень архитектуры системы.

2. Создание ИИ-агента выходит за рамки технологической задачи. Это вклад в технологический суверенитет, снижение зависимости от зарубежных решений и импульс для развития российской школы финансового ИИ.

3. Существующая нормативная база Банка России формирует лишь этические рамки и не содержит жестких требований к архитектуре ИИ-агента, что требует дальнейшей конкретизации.

4. Противодействие рискам требует не одного, а комплекса инструментов, таких как архитектурные решения против технических угроз,

институциональных сдержек для конфликтов интересов и поведенческих ограничений. Особую сложность представляют риски латентности рынка и поведенческой манипуляции, так как они находятся не в ошибках кода, а в природе рыночных взаимодействий и человеческой психологии.

5. Вопрос ответственности за ошибки ИИ-агента остается неурегулированным. Требуется четкое разделение функций и распределение ответственности.

6. Успех проекта определяется не технологической сложностью, а качеством институционального и архитектурного дизайна.

Рекомендации по управлению ИИ-агентом. На основе проведенного анализа рисков и с учетом предложенной нами модели распределения ответственности можно сформулировать следующие общие принципы управления ИИ-агентом.

1. Принцип человеческого контроля и приоритета человека.

Должна быть предусмотрена возможность переключения на обслуживание человеком-консультантом и пересмотра решений ИИ-агента по требованию клиента. Решение человека, принятое по запросу клиента, должно иметь приоритет над рекомендацией ИИ-агента.

2. Принцип полной аудируемости.

Все действия ИИ-агента (какие рекомендации, кому, на основании каких данных и критериев) должны фиксироваться в неизменяемом виде с возможностью последующего аудита.

3. Принцип ретроспективной прозрачности.

Полная спецификация алгоритма не должна раскрываться в режиме реального времени. Это не исключает «подгонку», то есть банки могут пытаться выявлять логику алгоритма эмпирически, но применение данного принципа будет затруднять эти попытки. Полная непрозрачность, однако, неприемлема: она делает невозможным общественный контроль за добросовестностью оператора. Для разрешения этого противоречия возможно применение механизма отложенного раскрытия спецификации каждой версии алгоритма.

Функция отложенного раскрытия – это обеспечение долгосрочной подотчетности. После публикации спецификации эксперты, СМИ и общественные организации могут проанализировать, не был ли алгоритм систематически смещен в пользу определенных банков в проверяемый период. Проблема же текущей «подгонки» продуктов под алгоритм выходит за рамки механизмов прозрачности и должна решаться комплексно, с применением иных мер: регулярной ротацией критериев ранжирования, запретом массового тестирования алгоритма банками, случайным аудитом и т. п.

4. Принцип запрета поведенческой манипуляции.

Интерфейс ИИ-агента, как представляется, должен быть нейтральным информационным инструментом, исключающим: геймифика-

цию (анимации, уведомления, индикаторы успеха, звуки, соревновательные механики); методы «подталкивания» (предустановленные опции, дефолтный выбор в пользу определенных продуктов, ограничение времени, визуальное выделение без объективных оснований); гиперперсонализацию, эксплуатирующую эмоциональные уязвимости пользователя, если она не служит его явным интересам. Возможна персонализация на основе объективных финансовых показателей пользователя (доходы, риск-профиль, инвестиционный горизонт), полученных с его согласия.

5. Принцип осторожной интерпретации данных для макропруденциального анализа.

Данные ИИ-агента могут использоваться Банком России для макропруденциального анализа только при соблюдении ряда условий: кросс-валидации, учета нерепрезентативности, противодействию манипуляции, предупреждении иллюзии точности, публичной отчетности.

6. Принцип распределения полномочий и взаимного контроля.

Реагирование и надзор должен осуществлять Банк России. Оператором может быть отдельная структура с целевым бюджетным финансированием. Возможно формирование общественного совета, осуществляющего внешний аудит конкурентной нейтральности (без дублирования функций регулятора и внутренней валидации). Функции разработки, валидации и эксплуатации алгоритма должны быть распределены между независимыми структурами.

7. Принцип финансовой ответственности и компенсации ущерба.

В рассматриваемой модели оператор не должен отвечать перед потребителем напрямую. Ответственность должны нести: банки-партнеры (за свои данные); банки – участники разработки (за дефекты на этапе создания, при условии, что будет доказана их вина в форме умысла или грубой неосторожности); компенсационный фонд, формируемый за счет обязательных отчислений банков-партнеров (за непредсказуемые ошибки). Оператор должен субсидиарно отвечать перед фондом при доказанной вине.

В более широком контексте практическая реализация рассматриваемой инициативы Банка России может стать важным прецедентом встраивания регуляторных требований непосредственно в алгоритмическую архитектуру финансового рынка – переход от регулирования правилами к регулированию кодом. Успешность этого перехода будет определяться не столько технологической сложностью, сколько способностью регулятора выстроить устойчивый баланс между автономией алгоритмов, фидуциарной ответственностью перед потребителями и четким распределением ответственности за возможный ущерб.

ЛИТЕРАТУРА / REFERENCES

1. Абрамова М.А., Дубова С.Е. Турбулентность угроз финансовой стабильности в новых реалиях развития денежной и платежной систем // Банковские услуги. 2022. № 7. С. 9–18. [Abramova M.A., Dubova S.E. Turbulence of threats to financial stability in the new realities of the development of monetary and payment systems // Banking Services. 2022. No. 7. Pp. 9–18. (In Russ.).] DOI: 10.36992/2075-1915_2022_7_9. EDN: ZRYZLM.
2. Абрамова М.А., Криворучко С.В., Луняков О.В., Фиापшев А.Б. Теоретико-методологический взгляд на предпосылки возникновения и особенности функционирования децентрализованных финансов // Финансы: теория и практика. 2025. Т. 29. № 1. С. 80–96. [Abramova M.A., Krivoruchko S.V., Lunyakov O.V., Fiapshev A.B. Theoretical and methodological perspective on the prerequisites of emergence and peculiarities of the functioning of decentralized finance // Finance: Theory and Practice. 2025. Vol. 29. No. 1. Pp. 80–96. (In Russ.).] DOI: 10.26794/2587-5671-2025-29-1-80-96. EDN: QICSEA.
3. Афанасьева О.Н., Сапожников Д.А. Эволюция научных представлений о финансовом посредничестве // Финансы и кредит. 2025. Т. 31. № 12. С. 35–52. [Afanas'yeva O.N., Sapozhnikov D.A. Evolution of scientific ideas about financial intermediation // Finance and Credit. 2025. Vol. 31. No. 12. Pp. 35–52. (In Russ.).] DOI: 10.24891/fc.31.12.35. EDN: LJQEHF.
4. Долматович И.А., Кешенкова Н.В. Нейронные сети: применение и перспективы использования в деятельности коммерческих банков // Сберегательное дело за рубежом. 2024. № 2. С. 41–52. [Dolmatovich I.A., Keshenkova N.V. Neural networks: application and prospects of use in the activities of commercial banks // Savings Business Abroad. 2024. No. 2. Pp. 41–52. (In Russ.).] DOI: 10.36992/2782-5949_2024_2_41. EDN: AELGQM.
5. Киселев А.С. О некоторых возможностях применения технологий искусственного интеллекта в банковской деятельности в целях обеспечения технологического суверенитета // Банковское право. 2026. № 1. С. 57–66. [Kiselev A.S. On some opportunities for applying artificial intelligence technologies in banking to ensure technological sovereignty // Banking Law. 2026. No. 1. Pp. 57–66. (In Russ.).] DOI: 10.18572/1812-3945-2026-1-57-66. EDN: CNGEQK.
6. Кочергин Д.А. Основные направления использования искусственного интеллекта в финансовой сфере // Вестник Института экономики Российской академии наук. 2025. № 6. С. 147–169. [Kochergin D.A. Main trends in the use of artificial intelligence in the financial sector // Bulletin of the Institute of Economics of the Russian Academy of Sciences. 2025. No. 6. Pp. 147–169. (In Russ.).] DOI: 10.52180/2073-6487_2025_6_147_169. EDN: EQJOKW.
7. Криничанский К.В., Анненская Н.Е. Понятие и перспективы финансового развития // Вопросы экономики. 2022. № 10. С. 20–36. [Krinichansky K.V., Annenskaya N.E. The concept and prospects of financial development // Voprosy Ekonomiki. 2022. No. 10. Pp. 20–36. (In Russ.).] DOI: 10.32609/0042-8736-2022-10-20-36. EDN: IKUSTF.
8. Симаева Е.П., Симаева Н.П. Правовые механизмы и поведенческая экономика: инструменты регулирования рынка сбережений // Экономика. Налоги. Право. 2025. Т. 18. № 4. С. 172–179. [Simaeva E.P., Simaeva N.P. Legal mechanisms and behavioral economics: instruments for regulating the savings market // Economics, Taxes & Law. 2025. Vol. 18. No. 4. Pp. 172–179. (In Russ.).] DOI: 10.26794/1999-849X-2025-18-4-172-179. EDN: PCOVGA.

9. Траченко М.Б., Стародубцева Е.Б., Кожанова А.В. Новая модель краудфандинга для защиты интересов инвесторов в России // Вестник Московского университета. Серия 6. Экономика. 2023. Т. 58. № 1. С. 45–62. [Trachenko M.B., Starodubtseva E.B., Kozhanova A.V. A new crowdfunding model to protect investors interests in Russia // Lomonosov Economics Journal. 2023. Vol. 58. No. 1. Pp. 45–62. (In Russ.).] DOI: 10.38050/MSU01300105-6-58-1-3. EDN: DQEZBT.
10. Хоминич И.П., Фрумина С.В. и др. Деньги, финансы, банки, страхование в цифровую эпоху: осмысление трансформаций, риски, рынки, финансовые институты / Под ред. И.П. Хоминич, С.В. Фруминой. М.: Русайнс, 2023. [Khomnich I.P., Frumina S.V. et al. Money, finance, banking, insurance in the digital era: understanding transformations, risks, markets, financial institutions / Ed. by I.P. Khomnich, S.V. Frumina. Moscow: Rusains, 2023. (In Russ.).] EDN: UMOTRJ.
11. Чердаков О.И. Отношение государства и бизнеса к разработке и внедрению технологий искусственного интеллекта в российскую финансовую сферу и экономику // Банковское право. 2022. № 5. С. 40–47. [Cherdakov O.I. The attitude of the state and business towards the development and implementation of artificial intelligence technologies in the Russian financial sector and economy // Banking Law. 2022. No. 5. Pp. 40–47. (In Russ.).] EDN: STQJFS.
12. Diamond D.W. Financial Intermediation and Delegated Monitoring // The Review of Economic Studies. 1984. No. 51 (3). Pp. 393–414.
13. Feng R., Li H., Liu M. Robo-Advisors Beyond Automation: Principles and Roadmap for AI-Driven Financial Planning // arXiv:2509.09922. 2025. <https://arxiv.org/abs/2509.09922> (accessed: 15.04.2026).
14. Hagiu A., Wright J. Multi-Sided Platforms // International Journal of Industrial Organization. 2015. Vol. 43. Pp. 162–174.
15. Inderst R., Ottaviani M. Misselling through Agents // American Economic Review. 2009. Vol. 99. No. 3. Pp. 883–908. DOI: 10.1257/aer.99.3.883.
16. Li S., Shen W. Confronting AI-Induced Power Disparity and Emotional Disruption: Expansive Duties of Broker-Dealers and Investment Advisers // Northwestern Journal of International Law & Business. 2024. Vol. 45. No. 1. P. 31.
17. Packin N.G., Kliger D., Reichman A., Rabinovitz S. Hooked and Hustled: The Predatory Allure of Gambified Finance // Columbia Business Law Review. 2025. Vol. 2025. No. 1.
18. Rochet J.C., Tirole J. Platform Competition in Two-Sided Markets // Journal of the European Economic Association. 2003. Vol. 1. No. 4. Pp. 990–1029.
19. Teplova T., Tomtosov A., Sokolova T. A Retail Investor in a Cobweb of Social Networks // PLOS ONE. 2022. Vol. 17. No. 12. P. e0276924. DOI: 10.1371/journal.pone.0276924.
20. Wang W. Regulating Algorithmic Accountability in Financial Advising // CLS Blue Sky Blog. 2025. June 4. <https://clsbluesky.law.columbia.edu/2025/06/04/regulating-algorithmic-accountability-in-financial-advising/> (accessed: 18.04.2026).

Дата поступления рукописи: 08.06.2026 г.

Дата принятия к публикации: 13.08.2026 г.

СВЕДЕНИЯ ОБ АВТОРАХ

Миленьков Александр Владимирович – доктор экономических наук, доцент, профессор кафедры мировых финансовых рынков и финтеха Российского экономического университета имени Г.В. Плеханова, Москва, Россия
ORCID: 0000-0001-7812-2348
milnikov.av@rea.ru

Фрумина Светлана Викторовна – кандидат экономических наук, доцент, заведующая кафедрой мировых финансовых рынков и финтеха Российского экономического университета имени Г.В. Плеханова; доцент кафедры общественных финансов Финансового университета при Правительстве Российской Федерации, Москва, Россия
ORCID: 0000-0001-5143-9417
frumina.sv@rea.ru

ABOUT THE AUTHORS

Aleksandr V. Milenkov – Dr. Sci. (Econ.), Associate Professor, Professor of the Department of Global Financial Markets and Fintech, Plekhanov Russian University of Economics, Moscow, Russia
ORCID: 0000-0001-7812-2348
milnikov.av@rea.ru

Svetlana V. Frumina – Cand. Sci. (Econ.), Associate Professor, Head of the Department of Global Financial Markets and Fintech, Plekhanov Russian University of Economics; Associate Professor of the Department of Public Finance, Financial University under the Government of the Russian Federation, Moscow, Russia
ORCID: 0000-0001-5143-9417
frumina.sv@rea.ru

AI-DRIVEN FINANCIAL ADVISORY: FROM OVERCOMING INFORMATION ASYMMETRY TO NEW RISKS AND REGULATORY CHALLENGES

The article analyzes the Bank of Russia's initiative to create a specialized AI agent for selecting financial products, which was announced in April 2026. It examines architectural and regulatory approaches to implementing agentic systems in financial advisory services. Based on a synthesis of international and Russian research, the study explores the economic nature of AI intermediation as a mechanism for overcoming information asymmetry, which simultaneously creates new challenges: technical vulnerabilities of algorithms, institutional risks, user interest manipulation, market latency, and the problem of liability distribution. Special attention is paid to the evolution of regulatory approaches abroad, the identification of factors determining the success of agentic systems implementation, as well as the analysis of international precedents. The author proposes a classification of maturity levels for AI financial intermediaries. Recommendations are formulated for the governance of the AI agent aimed at balancing technological autonomy, fiduciary responsibility, and competitive neutrality.

Keywords: *artificial intelligence, AI agent, financial advisory, information asymmetry, Bank of Russia, AI technical risks, behavioral risks, market latency, responsibility for AI recommendations.*

JEL: G18, G23, G28, O33.