

Возможности и риски применения искусственного интеллекта в деятельности центральных банков

Кочергин Дмитрий Анатольевич

Институт экономики РАН, Москва, Россия, e-mail: kda2001@gmail.com

Цитирование: Кочергин Д.А. (2026). Возможности и риски применения искусственного интеллекта в деятельности центральных банков. *Terra Economicus* 24(2), 101–116. DOI: 10.18522/2073-6606-2026-24-2-101-116

Статья посвящена исследованию возможностей и рисков, связанных с внедрением технологий искусственного интеллекта в центральных банках. В ходе исследования рассмотрены основные понятия и элементы иерархии искусственного интеллекта, определены наиболее перспективные направления его применения в деятельности центральных банков, выявлены риски, связанные с внедрением искусственного интеллекта, и предложены методы по их минимизации. В результате исследования установлено, что основными современными кейсами применения искусственного интеллекта в центральных банках являются: экономический анализ и прогнозирование, совершенствование функционирования платежных систем, регулирование, надзор и контроль за деятельностью финансовых организаций, обнаружение информационных аномалий и оценка рисков. Установлено, что использование технологий искусственного интеллекта в деятельности центральных банков способно повысить эффективность противодействия отмыванию денежных средств, укрепить кибербезопасность и содействовать реализации целевых мандатов в сфере ценовой и финансовой стабильности. Вместе с тем внедрение таких технологий может создавать новые источники операционного, репутационного и модельного рисков, а также рисков, связанных с конфиденциальностью данных, что требует формирования адекватной системы управления рисками. Адаптация существующих моделей управления рисками к проектам искусственного интеллекта на основе многоуровневой защиты, а также комплексное применение мер по минимизации негативных последствий внедрения искусственного интеллекта способно стимулировать его внедрение в различных сферах деятельности центральных банков.

Ключевые слова: искусственный интеллект; машинное обучение; генеративный искусственный интеллект; центральные банки; операции центрального банка; риски искусственного интеллекта

Benefits and risks of the application of artificial intelligence in central bank operations

Dmitry A. Kochergin

Institute of Economics RAS, Moscow, Russia, e-mail: kda2001@gmail.com

Citation: Kochergin D.A. (2026). Benefits and risks of the application of artificial intelligence in central bank operations. *Terra Economicus* 24(2), 101–116 (in Russian). DOI: 10.18522/2073-6606-2026-24-2-101-116

The article deals with the opportunities and risks associated with the implementation of artificial intelligence technologies in central banks. The study examines the key concepts and elements of the artificial intelligence hierarchy, identifies the most promising areas for the application of generative artificial intelligence in central bank operations, highlights the risks associated with the implementation of artificial intelligence, and proposes methods to minimize them. The study found that the main contemporary cases for artificial intelligence use in central banks involve: economic analysis and forecasting, improving the functioning of payment systems, regulation, supervision, and control over the activities of financial institutions, detection of information anomalies, and risk assessment. It has been established that the use of artificial intelligence technologies in central banking operations can enhance the effectiveness of anti-money laundering efforts, strengthen cybersecurity, and support the fulfillment of mandates related to price and financial stability. At the same time, the adoption of such technologies may create new sources of operational, reputational, and model risks, as well as risks related to data privacy, which requires the establishment of an adequate risk management system. Adapting existing risk management models to artificial intelligence projects based on a multi-layered defense approach, along with the comprehensive implementation of measures to minimize the negative consequences of artificial intelligence adoption, can facilitate its implementation across various areas of central bank operations.

Keywords: artificial intelligence; machine learning; generative artificial intelligence; central banks; central bank operations; artificial intelligence risks

JEL codes: C63, G21, G22, E58, O33

Введение

Информационные технологии с каждым годом играют все более заметную роль в экономическом развитии как развитых стран, так и юрисдикций с формирующимся рынком. Среди множества прорывных технологий, включая распределенные реестры (distributed ledger technology, DLT), блокчейн (blockchain), аналитику больших данных (big data analytics), интернет вещей (internet of things, IoT), облачные вычисления (cloud computing), виртуальную реальность (virtual reality, VR), квантовые вычисления (quantum computing) и технологии кибербезопасности (cybersecurity), именно с искусственным интеллектом (artificial intelligence, AI) связываются наибольшие ожидания ввиду широких возможностей и многообразия сфер его применения в экономике, включая финансовый сектор.

Применение искусственного интеллекта способно оказывать влияние не только на экономический рост и эффективность деятельности производственных компаний, финансовых организаций и денежно-кредитных регуляторов в странах, находящихся в технологическом авангарде, но и в странах догоняющего развития. Вместе с тем особенности моделей искусственного интеллекта, равно как их экономический потенциал и возможные направления применения, остаются недостаточно изученными вследствие стремительного развития данной технологии и появления новых кейсов использования различных AI-

моделей в финансовой сфере. Особый научный интерес представляют исследования возможностей применения искусственного интеллекта в деятельности центральных банков, которые выступают не только гарантами ценовой и финансовой стабильности, но и активными пользователями данных технологий.

Цель статьи состоит в определении новых возможностей и рисков, связанных с внедрением искусственного интеллекта в деятельность центральных банков. В ходе исследования рассмотрены основные понятия и элементы иерархии развития искусственного интеллекта, определены наиболее перспективные направления его применения в работе денежно-кредитных регуляторов, выявлены риски, сопутствующие внедрению искусственного интеллекта, а также предложены подходы к формированию системы управления рисками и меры по их минимизации в деятельности центральных банков.

Основные понятия и элементы иерархии развития искусственного интеллекта

Термин «искусственный интеллект» (artificial intelligence, AI) является широким понятием, обозначающим компьютерные системы, которые выполняют задачи, требующие интеллекта, подобного человеческому¹. Наилучшим образом современное развитие искусственного интеллекта можно охарактеризовать через его ключевые реперные точки, к которым относятся генеративный искусственный интеллект, все шире применяемый в повседневной и финансово-экономической деятельности, и общий искусственный интеллект, разработка которого остается долгосрочной научной целью. Генеративный искусственный интеллект (generative artificial intelligence, GenAI) представляет собой класс моделей искусственного интеллекта, способных создавать оригинальный контент, включая текст, изображения, видео, аудио и программный код, в ответ на запрос пользователя на естественном языке. В то время как общий (универсальный) искусственный интеллект (artificial general intelligence, AGI) представляет собой гипотетический уровень развития AI, характеризующийся наличием когнитивных способностей, сопоставимых с человеческими.

Истоки искусственного интеллекта были заложены в 1930–1950-е гг., когда были разработаны машина Тьюринга, основы алгоритмической теории и созданы первые электронные вычислительные машины. Дальнейшее развитие машинного обучения (machine learning, ML)² и нейронных сетей (neural networks, NN)³ в 1950–1990-е гг., наряду со стремительным ростом вычислительных мощностей, обеспечило значительный прогресс в области искусственного интеллекта и расширило возможности автоматизации интеллектуальных задач. В 1980–2010-е гг. развитие нейронных сетей привело к появлению многослойных архитектур с большой глубиной (network's depth)⁴, что стало основой для реализации прикладных решений, включая системы распознавания лиц и голосовых помощников, получивших широкое распространение в финансовом секторе. Последующее развитие глубокого обучения (deep learning, DL)⁵ на основе современных нейронных сетей существенно расширило возможности обработки сложных данных и способствовало развитию методов обработки естественного языка (natural language processing, NLP)⁶, лежащих в основе современных моделей искусственного интеллекта⁷.

Появление NLP позволило разработать модели искусственного интеллекта, которые первоначально получили широкое применение в машинном переводе, поиске и ранжировании информации, классификации документов и автоматическом резюмировании, а впоследствии – в задачах обработки неструктурированных данных, включая картирование пространства экономических идей, выявление скрытых точек зрения и методологических подходов, лежащих в основе принятия решений по ключевым экономическим параметрам – таким как ключевая ставка и валютный

¹ Bank for International Settlements (2024). Artificial intelligence and the economy: Implications for central banks. BIS annual economic report, June, pp. 91–127. <http://www.bis.org/publ/arpdf/ar2024e3.pdf> (accessed on May 10, 2026)

² Машинное обучение (ML) – это собирательный термин, обозначающий методы, предназначенные для обнаружения закономерностей в данных и их использования для прогнозирования или принятия решений.

³ Нейронная сеть (NN) – это математическая модель, предназначенная для распознавания закономерностей и обучения на данных, которая принимает решения аналогично человеческому мозгу, используя процессы, имитирующие работу биологических нейронов.

⁴ Глубина сети – это количество слоев, которые позволяют нейронным сетям улавливать сложные взаимосвязи в данных. Глубокие сети с большим количеством параметров требуют наличия больших данных для обучения, но и предсказывают точнее.

⁵ Глубокое обучение (DL) – это разновидность методов машинного обучения, основанная на использовании многослойных нейронных сетей для построения многоуровневых представлений данных и выявления скрытых закономерностей.

⁶ Обработка естественного языка (NLP) – область применения ML для работы с языком, позволяющая машинам понимать лингвистику, определять значение контекста и самим генерировать смысловые фразы и тексты.

⁷ См.: Russell, Norvig, 2021.

курс⁸. Впоследствии в рамках задач NLP была разработана архитектура «последовательность в последовательность» (sequence-to-sequence, seq2seq)⁹, которая стала широко применяться для прогнозирования временных рядов (цен на акции, объемов торгов активами и др.), обработки финансовых документов и автоматической генерации описаний транзакций¹⁰.

Логическим продолжением развития методов обработки естественного языка (NLP) стало появление в конце 2010-х гг. моделей искусственного интеллекта на основе трансформеров (transformer models). Модели-трансформеры предназначены для обработки последовательностей с использованием механизма многоголового внимания (multi-head attention). Благодаря параллельной обработке последовательностей, эффективной работе с длинными зависимостями, высокой масштабируемости и универсальности трансформеры легли в основу создания больших языковых моделей. Большие языковые модели (large language models, LLMs)¹¹ стали одной из наиболее продвинутых категорий моделей в рамках NLP и сыграли важную роль в развитии искусственного интеллекта. Их появление позволило учитывать широкий контекст входных данных, улучшить улавливание смысловых нюансов текста, включая задачи перевода и анализа тональности¹². Большие языковые модели способны значительно точнее, чем их предшественники, интерпретировать смысловые зависимости и генерировать релевантные ответы при минимальном количестве примеров или их полном отсутствии. Их применение позволило широкому кругу пользователей с помощью естественного языка автоматизировать задачи, ранее выполнявшиеся специализированными моделями, требующими профессиональных навыков.

На базе больших языковых моделей в конце 2010-х – начале 2020-х гг. сформировалось направление генеративного искусственного интеллекта (generative artificial intelligence, GenAI). Хотя понятия LLMs и GenAI часто используются в близком контексте, GenAI является более широким понятием, в то время как LLMs выступают одним из его ключевых компонентов. Если LLMs ориентированы преимущественно на работу с текстовыми данными, то GenAI охватывает широкий класс систем, предназначенных для генерации нового контента, включая текст, изображения, аудио и видео. Генеративные модели способны, например, анализировать визуальные данные, учитывать сопутствующий текстовый контекст и генерировать новые изображения с заданными характеристиками. Так, подобные возможности реализованы в системах GenAI, к числу которых относятся ChatGPT, DALL-E, Sora и Claude.

В настоящее время применение GenAI в финансовом секторе активно расширяется¹³. Такие системы используются для генерации финансовых отчетов, подготовки рыночных обзоров, разработки торговых стратегий, анализа настроений участников рынка, формулирования гипотез, автоматизации комплаенса и нормативной отчетности, а также создания AI-ассистентов. Среди наиболее известных примеров прикладного использования GenAI в финансах можно выделить модель BloombergGPT и прикладное AI-решение JPMorgan IndexGPT.

Новейшим этапом развития технологий искусственного интеллекта (эволюционным развитием GenAI) являются AI-агенты (AI agents)¹⁴, представляющие собой автономные системы, ориентированные на достижение заданных целей (например, максимизацию доходности при заданном уровне риска или диверсификацию портфеля). Такие системы способны к планированию и поэтапному выполнению действий, включая проведение сделок, подготовку отчетов и формирование аналитических обзоров рынков, а также к обучению и адаптации на основе новых данных и результатов. В отличие от генеративных моделей искусственного интеллекта, AI-агенты обладают более высоким уровнем автономности, позволяющим им самостоятельно адаптироваться к изменениям внешней среды и корректировать стратегии поведения. Современными примерами подоб-

⁸ Bank for International Settlements, 2024, указ. соч.

⁹ «Последовательность в последовательность» (seq2seq) – это нейросетевая архитектура, применяемая в задачах NLP для преобразования одной текстовой последовательности в другую, особенно эффективная в случаях, когда длины входной и выходной последовательностей различаются.

¹⁰ Подробнее об эволюции технологий AI см.: Гаврилова, 2024; Кочергин, 2025.

¹¹ Большие языковые модели (LLMs) – это модели AI, обученные на большом количестве параметров или объемах данных, что делает их способными понимать и генерировать естественный язык и другие типы контента для выполнения широкого спектра задач.

¹² Подробнее см.: Park et al., 2025.

¹³ Подробнее о направлениях внедрения AI в финансовой сфере см.: Банк России (2023). Применение искусственного интеллекта на финансовом рынке. Доклад для общественных консультаций. https://cbr.ru/Content/Document/File/156061/Consultation_Paper_03112023.pdf; Кочергин, 2025.

¹⁴ AI-агенты – это автономные системы, основанные на GenAI, способные воспринимать окружающую среду, анализировать ее, принимать решения и выполнять действия для достижения заданной цели.

ных систем являются AutoGPT¹⁵, Devin (Cognition AI)¹⁶ и агенты на базе GPT-моделей. К наиболее перспективным направлениям применения AI-агентов в финансовом секторе относятся алгоритмическая торговля и арбитраж, финансовое моделирование и прогнозирование, управление инвестиционными портфелями, а также обнаружение мошенничества и выявление киберугроз.

На рисунке 1 мы визуализировали эволюционно-иерархическую структуру технологий искусственного интеллекта.



Рис. 1. Эволюционно-иерархическая структура технологий искусственного интеллекта
Источник: составлено автором; Кочергин, 2025: 154.

Несмотря на значительный прогресс технологий искусственного интеллекта в последние десятилетия, сохраняется существенная неопределенность относительно долгосрочных возможностей LLMs и GenAI. Хотя современные модели способны выполнять задачи, требующие ограниченных когнитивных функций, и демонстрируют отдельные эмерджентные способности, их эффективность в задачах, требующих сложного многошагового логического рассуждения и устойчивого понимания, остается ограниченной (Perez-Cruz, Shin, 2024). Вероятно, данные ограничения связаны не столько с объемом обучающих данных или числом параметров моделей, сколько с преимущественно языковой природой их обучения, которая не включает непосредственного взаимодействия с физическим миром и, соответственно, не формирует человеческое знание и опыт, возникающие в результате такого взаимодействия (McCaul, 2024).

¹⁵ AutoGPT – первый массовый AI-агент с открытым исходным кодом на базе GPT, который самостоятельно разбивает задачи на последовательные шаги и выполняет их, руководствуясь поиском оптимального решения.

¹⁶ Devin (Cognition AI) – один из первых AI-разработчик, который сам пишет код, вводит в эксплуатацию и тестирует, выполняя функции виртуального инженера.

Основные кейсы применения AI в деятельности центральных банков

Влияние искусственного интеллекта на экономическую деятельность существенно варьируется между странами и отраслевыми секторами. Наибольшие выгоды от внедрения AI получают отрасли, основанные на когнитивно и информационно интенсивных видах деятельности, а также страны с высоким уровнем технологического развития, развитой цифровой инфраструктурой, качественным человеческим капиталом и благоприятной институциональной средой (Gambacorta et al., 2025). К таким сферам относится и деятельность центральных банков (ЦБ), характеризующаяся высокой интенсивностью обработки данных, аналитической направленностью и существенной зависимостью качества решений от информационного обеспечения.

В этом контексте потенциал применения искусственного интеллекта в ЦБ отличается значительным разнообразием. Кроме того, AI часто может применяться в связке с другими новыми информационными технологиями, такими как аналитика больших данных и распределенные реестры¹⁷. Согласно современным исследованиям, большинство ЦБ находятся на стадиях планирования, экспериментирования и пилотного внедрения AI-решений. При этом предполагается, что их институциональное и аналитическое преимущество в перспективе будет определяться не столько объемом накопленных данных, сколько способностью оперативно, безопасно и в стандартизированной форме преобразовывать их в аналитическое знание, необходимое для повышения качества решений в области денежно-кредитной политики, банковского надзора и обеспечения финансовой стабильности¹⁸ (Araujo et al., 2026).

В настоящее время использование AI-моделей в центральных банках преимущественно связано с анализом больших неструктурированных данных (тексты, отчеты, социальные медиа), что позволяет совершенствовать надзор и аналитику (включая выявление рисков и автоматизацию отдельных задач), усиливать кибербезопасность (обнаружение уязвимостей и аномалий), а также повышать эффективность внутренних операционных процессов (Araujo et al., 2025). При этом направления применения AI в центральных банках продолжают расширяться. Наиболее перспективными кейсами применения искусственного интеллекта в ЦБ являются:

1) *экономический анализ и прогнозирование*. Модели AI могут использоваться для прогнозирования макроэкономических показателей, таких как ВВП и инфляция¹⁹, а также для анализа влияния денежно-кредитной политики на финансовый сектор и макроэкономику;

2) *платежные системы*. AI может применяться для совершенствования функционирования платежной инфраструктуры²⁰, включая выявление аномальных транзакций, обнаружение сетей отмывания денег, а также анализ платежных систем с целью разработки стандартов и инноваций, повышающих их устойчивость и надежность;

3) *регулирование*. Модели обработки естественного языка и AI могут использоваться для оценки сложности пруденциального и банковского регулирования на основе алгоритмов классификации и анализа сходства, а также для анализа влияния нормативных актов и законопроектных на кредитные организации и финансовые рынки;

4) *надзор и контроль*. ML, NLP и GenAI могут применяться для поддержки надзорной деятельности, включая извлечение и предоставление информации надзорным органам, мониторинг финансовых институтов и оценку их кредитного и портфельного риска²¹;

5) *обнаружение аномалий данных*. Модели AI могут использоваться для повышения качества данных, включая выявление аномалий и ошибок в отчетности, искажений в ценовых данных и информации об остатках на счетах, а также для усиления кибербезопасности²² (Araujo et al., 2022);

¹⁷ См.: Болонин и др., 2024.

¹⁸ Подробнее см.: Bank for International Settlements (2025). The use of artificial intelligence for policy purposes. Report submitted to the G20 finance ministers and central bank governors, October <https://www.bis.org/publ/othp100.pdf> (accessed on May 10, 2026)

¹⁹ См.: Benford, J. (2024). Trusted AI: Ethical, safe and effective application of artificial intelligence at the Bank of England. Speech given at the Central Bank AI Conference, Bank of England, September. <https://www.bankofengland.co.uk/speech/2024/september/james-benford-speech-at-the-central-bank-ai-inaugural-conference> (accessed on May 10, 2026); Dauphin et al., 2022.

²⁰ Подробнее см.: BIS Innovation Hub (2023). Project Aurora: The power of data, technology and collaboration to combat money laundering across institutions and borders. May. <https://www.bis.org/publ/othp66.pdf> (accessed on May 10, 2026)

²¹ Bank for International Settlements, 2024, указ. соч.; Подробнее см.: Benford, 2024, указ. соч.; McCaul E. (2024). From Data to Decisions: AI and Supervision. ECB Supervision Newsletter. 26 February. <https://www.bankingsupervision.europa.eu/press/interviews/date/2024/html/ssm.in240226~c6f7fc9251.en.html> (accessed on May 10, 2026)

²² В ходе опроса, проведенного экономистами Банка международных расчетов (БМР), большинство экспертов ЦБ считает, что искусственный интеллект дает больше преимуществ, чем рисков, и что он может превзойти традиционные методы повышения эффективности в управлении кибербезопасностью (Aldasoro et al., 2024).

6) *оценка рисков*. Модели AI могут применяться для построения систем раннего предупреждения, выявления ранних признаков потенциальных рисков, анализа рисков событий и проведения стресс-тестирования²³;

7) *клиентские и корпоративные услуги*. AI может использоваться для аналитических и консультационных задач, предоставления общественных и корпоративных сервисов через чат-боты²⁴, а также для поддержки сотрудников в вопросах кадровой политики и развития информационных технологий, включая автоматизацию разработки кода и управления базами данных²⁵.

На рисунке 2 представлена градация основных возможностей применения технологий AI в деятельности ЦБ²⁶.

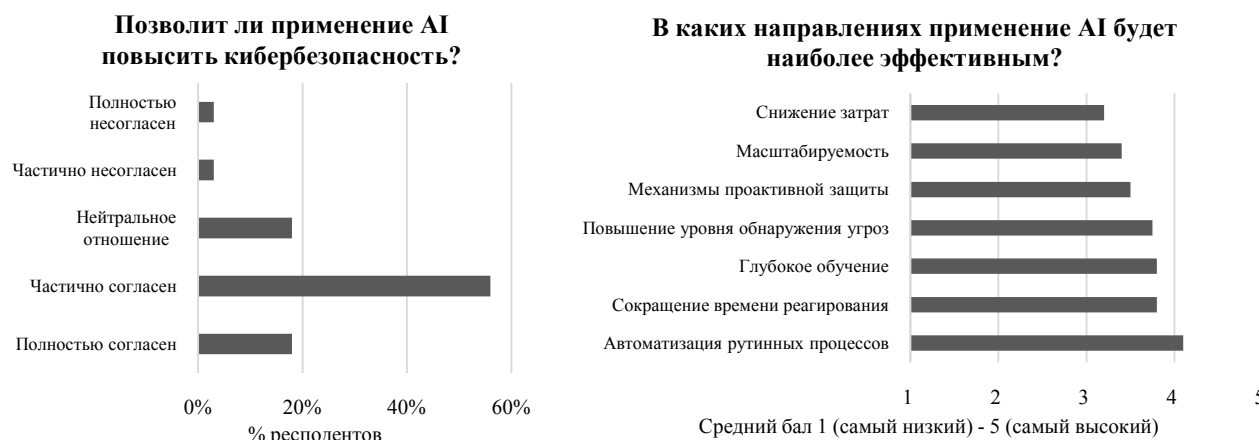


Рис. 2. Возможности AI для решения различных задач в ЦБ

Источник: Aldasoro et al., 2024: 10.

Как видно из рис. 2, по мнению представителей ведущих ЦБ, преимущества применения искусственного интеллекта связаны прежде всего с автоматизацией рутинных задач, позволяющей снизить затраты на трудоемкие виды деятельности, традиционно выполняемые человеком. Дополнительные преимущества использования искусственного интеллекта в киберпространстве включают повышение эффективности обнаружения угроз, сокращение времени реагирования на кибератаки, а также выявление новых тенденций, аномалий и корреляций, которые могут быть неочевидны для аналитиков-людей.

Широкое внедрение искусственного интеллекта может повлиять на уровень производительности, потребления и инвестиций, а также на структуру рынка труда²⁷, что, в свою очередь, способно воздействовать на уровень цен, эффективность и бесперебойность функционирования финансовой системы. В этой связи ЦБ следует более активно использовать модели искусственного интеллекта для поддержания ценовой и финансовой стабильности, совершенствования инструментов прогнозирования, а также реализации денежно-кредитной политики и надзорных функций. Например, ЕЦБ уже разрабатывает более сорока кейсов применения искусственного интеллекта в банковском надзоре²⁸. Бундесбанк использует NLP для автоматизированной поддержки решений по одобрению проспектов ценных бумаг, а также AI-агента «MILA» для анализа текстов монетарной политики с целью оценки тональности заявлений (Geiger et al., 2025). Федеральная резервная система США применяет GenAI для анализа протоколов заседаний (Baily et al., 2025; Dunn et al.,

²³ Подробнее см.: Nistor, M. (2023). Artificial intelligence applications in finance: Considering the benefits and risks. The Teller Window. Federal Reserve Bank of New York, December 14. <https://tellerwindow.newyorkfed.org/2023/12/14/ai-applications-in-finance-considering-the-benefits-and-risks/> (accessed on May 10, 2026)

²⁴ См.: Benford, 2024, указ. соч.

²⁵ Bank for International Settlements, 2024, указ. соч.

²⁶ По данным опроса БМР, около 71% центральных банков используют GenAI для повышения кибербезопасности (Aldasoro et al., 2024).

²⁷ Более 83% ЦБ отмечают, что кадровое планирование стало более сложным из-за внедрения AI (Bell et al., 2025). Хотя в этой функции AI пока преимущественно применяется как «помощник» (co-pilot), но в перспективе он может стать «заменителем» (agent) людей.

²⁸ Hoppe, F. (2025). Benefits from advanced technology infrastructure in supervision. ECB Supervision Newsletter, 14 May. <https://www.bankingsupervision.europa.eu/presssupervisory-newsletters/newsletter/2025/html/ssm.nl250514.en.html> (accessed on May 10, 2026); OECD, 2026.

2024), а центральные банки Канады и Японии используют AI-модели для повышения эффективности экономического анализа, связанного с деятельностью финансовых институтов^{29, 30}.

Совместные проекты по использованию AI в ЦБ

В настоящее время центральные банки и финансовые учреждения различных стран под эгидой инновационных хабов Банка международных расчетов координируют усилия по реализации совместных пилотных проектов, направленных на изучение потенциала использования технологий искусственного интеллекта в различных областях деятельности ЦБ и частных финансовых учреждений. Например, в проекте Raven, реализуемом под управлением Скандинавского инновационного хаба БМР, технологии AI применяются для повышения устойчивости к киберугрозам. Проект показал, что искусственный интеллект может существенно повысить эффективность выявления и анализа киберугроз в финансовой инфраструктуре за счет автоматизированного распознавания аномалий и ускоренной обработки больших массивов событийных данных. Кроме того, было продемонстрировано, что использование AI позволяет повысить устойчивость финансовых систем к киберинцидентам, сокращая время обнаружения угроз и улучшая качество реагирования³¹.

В рамках проекта Neo, реализуемого под эгидой Швейцарского инновационного хаба Банка международных расчетов, разрабатывается модель прогнозирования основных макроэкономических показателей на основе постоянно обновляющихся данных о состоянии различных секторов экономики. Результаты проекта Neo показали, что методы машинного обучения позволяют ЦБ повышать оперативность и точность оценки текущей экономической активности, а также краткосрочных макроэкономических прогнозов, особенно в условиях задержек традиционной статистики.

В проекте Spectrum, реализуемом под управлением Евросистемного хаба БМР, разрабатывается модель структурирования больших массивов микроданных по ценам широкой товарной номенклатуры для прогнозирования инфляции. Результаты проекта продемонстрировали, что использование искусственного интеллекта для преобразования неструктурированных микроценовых данных в унифицированные товарные категории позволяет существенно ускорить и масштабировать процесс инфляционного прогнозирования без потери точности³². Это подтверждает потенциал применения искусственного интеллекта для оперативного мониторинга инфляционных процессов и совершенствования аналитической поддержки монетарной политики³³.

В рамках проекта Noor, реализуемого Инновационным хабом Банка международных расчетов совместно с Денежно-кредитным управлением Гонконга и Управление по финансовому регулированию и надзору Великобритании, исследуются возможности применения AI в надзорной деятельности ЦБ посредством методов объяснимого искусственного интеллекта (explainable artificial intelligence, XAI)³⁴. Данные методы используются для анализа алгоритмических моделей искусственного интеллекта, применяемых финансовыми организациями. Результаты проекта подтвердили возможность использования искусственного интеллекта для повышения прозрачности, оценки устойчивости и независимой верификации решений алгоритмов в целях риск-ориентированного финансового надзора³⁵.

²⁹ Bank of Canada (2024). Canadians count on us: The Bank of Canada's 2025–27 strategic plan. <https://www.bankofcanada.ca/about/governance-documents/the-bank-of-canadas-2025-27-strategic-plan/the-bank-of-canadas-2025-27-strategic-plan-framework/#framework> (accessed on May 10, 2026)

³⁰ Bank of Japan (2025). Use and risk management of generative AI by Japanese financial institutions. The results of FY2025 Survey. Financial system report annex series, September. <https://www.boj.or.jp/en/research/brp/fsr/data/fsrb250930.pdf> (accessed on May 10, 2026)

³¹ BIS Innovation Hub (2024). Project Raven: Using AI to assess financial system's cyber security and resilience, April. http://www.bis.org/about/bisih/topics/cyber_security/raven.htm (accessed on May 10, 2026)

³² BIS Innovation Hub (2026). Project Spectrum: Using generative AI to enhance inflation nowcasting, February. <https://www.bis.org/publ/othp109.htm> (accessed on May 10, 2026)

³³ В проекте Raven используются технологии искусственного интеллекта в области NLP, включая LLMs. В проекте Neo задействованы LLMs, а также методы «обучения с учителем» (supervised machine learning, SML), при котором модель обучается на размеченных данных с заранее известными правильными ответами, и «обучения без учителя» (unsupervised machine learning, UML), при котором модель работает с неразмеченными данными без заранее заданных правильных ответов. В проекте Spectrum применяются технологии искусственного интеллекта в области NLP, включая методы UML и LLMs (BIS, 2024, указ. соч.).

³⁴ Методы обеспечения объяснимости искусственного интеллекта (XAI) – это совокупность методов и подходов, позволяющих объяснять логику работы моделей искусственного интеллекта и причины принимаемых ими решений.

³⁵ BIS Innovation Hub (2025). Project Noor: Explaining AI models for financial supervision, August. https://www.bis.org/about/bisih/topics/suptech_regtech/noor.htm (accessed on May 10, 2026)

Среди наиболее известных проектов Банка международных расчетов по изучению возможностей применения искусственного интеллекта в платежной и расчетной сфере можно выделить Agorá и Ауорога. В рамках проекта Agorá, объединяющего семь центральных банков и более 40 финансовых учреждений частного сектора, создается единая программируемая платформа для унификации расчетов по операциям с токенизированными активами финансовых организаций, при которых урегулирование осуществляется с использованием цифровых валют ЦБ (CBDC). Платформа Agorá объединяет обмен сообщениями и обработку платежей в рамках одной операции, позволяет проводить расчеты атомарно и обеспечивает выполнение процедур KYC / AML с сохранением конфиденциальности пользователей³⁶. В отличие от корреспондентского банкинга, где проверка информации и обновление счетов осуществляются раздельно, что приводит к дублированию операций, на платформе Agorá обмен финансовыми сообщениями, движение активов и клиринг интегрированы в единый атомарный процесс, устраняющий риск аннулирования транзакций и возврата средств³⁷. Применение искусственного интеллекта в данном проекте не только повышает уровень конфиденциальности и упрощает процессы комплаенса в рамках KYC/AML, но и позволяет автоматизировать контроль соблюдения регуляторных требований на этапе предварительного отбора транзакций.

В рамках проекта Ауорога, реализуемого под патронажем Скандинавского инновационного хаба Банка международных расчетов, исследуются возможности применения искусственного интеллекта для выявления случаев отмывания денег и противодействия финансовым преступлениям на трансграничном уровне. Посредством имитации схем отмывания денег в рамках проекта Ауорога изучаются возможности применения различных моделей машинного обучения в платежных операциях. Сравнение проводится по следующим сценариям: данные о транзакциях изолированы на уровне отдельного банка; данные о транзакциях агрегированы на национальном уровне; платежные данные объединены на трансграничном уровне. Модели обучаются на синтетических данных, имитирующих отмывочные операции, и впоследствии используются для оценки вероятности отмывания денег в каждом из сценариев³⁸. В результате реализации проекта были получены два ключевых вывода. Во-первых, модели машинного обучения превосходят традиционные методы выявления отмывания денег, используемые в большинстве юрисдикций. Во-вторых, их эффективность особенно возрастает при объединении данных из разных финансовых учреждений в одной или нескольких юрисдикциях, что подчеркивает важность трансграничной координации в сфере противодействия отмыванию денег (AML)³⁹.

Типологизация рисков внедрения AI в ЦБ

Несмотря на отмеченные выше преимущества применения искусственного интеллекта и успешное пилотирование соответствующих моделей как на уровне отдельных ЦБ, так и в рамках совместных проектов, внедрение данной технологии, как и любой другой новой информационной технологии, в деятельность денежно-кредитных регуляторов неизбежно связано как с традиционными, так и с новыми видами рисков. В таблице систематизированы и описаны основные риски, с которыми могут столкнуться центральные банки при внедрении искусственного интеллекта.

Как следует из таблицы, применение искусственного интеллекта в деятельности ЦБ, помимо общих рисков, связанных с внедрением новых информационных технологий, приводит к возникновению нового типа риска – риска моделей AI. Данный риск связан с возможностью неверной интерпретации и недостаточной надежностью генерируемой информации, а также с недостаточной прозрачностью и подотчетностью разработчиков моделей искусственного интеллекта. Так, низкое качество данных и некорректные подсказки могут приводить к ошибочным анализам и прогнозам, а также к воспроизведению предвзятостей и систематических ошибок.

³⁶ См.: Private sector partners join project Agorá. Bank for International Settlements, 2025. <https://www.bis.org/about/bisih/topics/fmis/agora.htm> (accessed on May 10, 2026)

³⁷ Подробнее см.: Garratt et al., 2024.

³⁸ BIS Innovation Hub (2025). Project Aurora: The power of data, technology and collaboration to combat money laundering across institutions and borders. <https://www.bis.org/about/bisih/topics/fmis/aurora.htm> (accessed on May 10, 2026)

³⁹ В данном проекте применяются графовые нейронные сети, являющиеся специальными искусственными нейронными сетями, предназначенными для решения задач, входными данными для которых являются графы, которые продемонстрировали отличную производительность при выявлении подозрительных платежных транзакций.

Таблица

Классификация рисков внедрения AI в ЦБ

Вид риска	Описание риска	Применим к новым информац. технологиям в целом	Применим непосредственно к AI
Стратегический риск	Отсутствие четкой стратегии использования искусственного интеллекта, а также ясных принципов его управления может негативно сказаться на репутации ЦБ через возникновение и закрепление смещений (bias) в алгоритмах, что в свою очередь может препятствовать достижению его целей	X	X
Операционные риски	Внедрение искусственного интеллекта может осложнить деятельность ЦБ, усиливая ряд операционных рисков, включая <i>риск правовой неопределенности</i> (несоответствие результатов работы AI юридическим обязательствам), <i>комплаенс-риск</i> (недостаточная объяснимость моделей, затрудняющая подтверждение соответствия нормативным требованиям), <i>риски, связанные с человеческим фактором</i> (недостаток компетенций и подготовки сотрудников в области безопасного применения AI, что может создать стратегическое отставание), а также <i>риск нарушения непрерывности бизнеса</i> (усиление традиционных угроз, включая незаконную деятельность, терроризм, военные конфликты и др., осложняющих реагирование на сбои). Указанные риски могут приводить к убыткам вследствие неэффективности внутренних процессов, ошибок персонала, сбоев систем или внешних факторов	X	X
Риски конфиденциальности и кибербезопасности	Результаты работы AI могут содержать конфиденциальную информацию, которая при неправомерном использовании способна привести к нарушению безопасности пользователей, несанкционированному доступу к финансовым данным и вмешательству в частную жизнь. Это повышает риск утечек и искажения персональных данных, нарушения прав интеллектуальной собственности и совершения иных финансовых преступлений. Интеграция AI-систем с инфраструктурой ЦБ создает дополнительные векторы кибератак, включая инъекции подсказок, отравление обучающих данных и кражу моделей, что может привести к нарушению целостности, достоверности и полноты данных	X	X
Риски информационно-коммуникационных технологий (ИКТ)	Функционирование AI может сопровождаться ошибками программного обеспечения, сбоями оборудования и неадекватным дизайном AI-систем, что может нарушать непрерывность бизнеса, снижать производительность, доступность и интеграцию систем, а также ставить под угрозу персональные данные пользователей. Дополнительным источником риска являются проблемы совместимости новых технологий с существующей ИКТ-инфраструктурой, которые могут усиливать сбои, снижать надежность систем и негативно влиять на эффективность финансовых операций	X	X

Окончание табл.

Вид риска	Описание риска	Применим к новым информац. технологиям в целом	Применим непосредственно к AI
Риск третьей стороны	Зависимость центральных банков от внешних AI-моделей и инструментов, предоставляемых сторонними поставщиками, формирует риск конфиденциальности и информационной безопасности (неправомерное использование данных), риск концентрации (влияние сбоя у поставщика на деятельность ЦБ), а также риск кибербезопасности (воздействие кибератак или технических сбоев на стороне провайдера на функционирование систем ЦБ)	X	X
Экологические, этические и социальные риски	Риски связаны с ростом энергопотребления и углеродных выбросов вследствие широкого использования энергоемких вычислительных ресурсов для обучения и обработки больших данных. Дополнительно риски обусловлены возможной генерацией AI неадекватных с этической и социальной точки зрения ответов из-за ограниченности понимания системами искусственного интеллекта социального и этического контекста человеческой деятельности	X	X
Репутационные риски	Могут возникать как следствие реализации любых из перечисленных выше рисков и выражаться в формировании негативного общественного мнения из-за ошибок, утечек данных, а также недостаточной прозрачности и контроля над AI-моделями, что способно привести к снижению доверия к ЦБ	X	X
Риски моделей AI	Связаны с ограниченной интерпретируемостью и надежностью результатов работы искусственного интеллекта вследствие недостаточной прозрачности и подотчетности моделей, их склонности к «галлюцинациям», а также потенциальной возможности неконтролируемой модификации алгоритмов		X

Источник: составлено автором по: BIS, 2025: 7–9.

Модели AI, как отмечалось ранее, также подвержены ошибкам рассуждений и «галлюцинациям». Аналогичным образом использование AI при работе с большими данными может приводить к апофении – выявлению закономерностей там, где они фактически отсутствуют. Это связано с тем, что большие объемы данных могут порождать множественные, часто случайные взаимосвязи (Boyd, Crawford, 2012). Недостаток публичной и детализированной информации об особенностях используемых AI-моделей может создавать проблемы управления (в частности, недостаточный контроль за данными, используемыми для обучения алгоритмов) и снижать подотчетность разработчиков и операторов таких систем (BIS, 2025).

Помимо рисков, при внедрении искусственного интеллекта в ЦБ следует учитывать и иные направления его влияния, которые вследствие сложного и динамичного характера данных технологий могут приводить к неожиданным результатам. К таким направлениям относятся: незапланированное использование (AI может не только выявлять новые проблемы в деятельности финансовых учреждений, но и способствовать их возникновению); влияние на сотрудников (неравномерное распределение выгод от применения AI в различных должностных ролях); сложности внедрения (недооценка затрат на обучение, развертывание моделей и их интеграцию с существующими системами); а также зависимость от AI (формирование технологической зависимости информационных процессов от AI-инструментов, сопровождаемое ростом расходов на их модернизацию в финансовых организациях и ЦБ) (BIS, 2025).

Принципы регулирования AI и управление их рисками в ЦБ

В связи с наличием рисков и дополнительных направлений влияния, связанных с внедрением искусственного интеллекта, ряд международных организаций⁴⁰ в последние годы разработали принципы управления искусственным интеллектом для регуляторов, организаций и иных заинтересованных сторон⁴¹. Применительно к центральным банкам данные принципы направлены на обеспечение низкой склонности к риску, поддержку выполнения мандатов и прозрачности деятельности⁴². Для формирования доверия к системам искусственного интеллекта они должны соответствовать следующим принципам: 1) безопасность (надежность, защищенность и устойчивость систем); 2) конфиденциальность (внутреннее управление и защита данных); 3) объяснимость и прозрачность принимаемых решений⁴³; 4) надежность (стабильное выполнение функций в соответствии с ожиданиями); 5) ответственность (сохранение человеческого контроля над решениями AI); 6) этичность (недискриминационность, справедливость и вклад в общественное благополучие) и др.⁴⁴ Реализация указанных принципов требует создания адаптивных систем управления рисками внедрения искусственного интеллекта в ЦБ.

В основе системы управления рисками искусственного интеллекта в центральных банках должна лежать комплексная стратегия, определяющая профиль рисков AI, обеспечивающая отбор и оценку проектов с учетом этих рисков, адаптацию существующих моделей риск-менеджмента к AI-проектам, а также защиту информации на всех этапах жизненного цикла. Формирование профиля рисков позволяет ЦБ оценивать соответствие результатов внедрения AI их риск-аппетиту и возможность согласования нормативных требований с целями в области финансовой устойчивости⁴⁵. На начальных этапах внедрения AI целесообразно ограничивать его использование внутренними процессами, не относящимися к критически важным функциям, до тех пор, пока связанные риски не будут надлежащим образом оценены и контролируемы.

Оценка и отбор проектов в области искусственного интеллекта должны осуществляться в соответствии со стратегическими и инновационными целями центрального банка и с учетом допустимого уровня риска, что позволяет определить модели и инструменты искусственного интеллекта, соответствующие риск-аппетиту ЦБ. На первом этапе целесообразно выявить и проанализировать случаи использования AI, а также классифицировать их по уровню риска в зависимости от чувствительности данных, используемых для обучения. На втором этапе следует определить надлежащие меры контроля, обеспечивающие баланс между рисками и их управлением в соответствии с риск-профилем ЦБ. На третьем этапе, после принятия решения о целесообразности реализации проекта AI, необходимо разработать и внедрить соответствующие контрольные меры, а также адаптировать под них существующую технологическую инфраструктуру. Проекты AI должны быть интегрированы в действующие механизмы управления рисками и подлежать мониторингу на протяжении всего жизненного цикла (BIS, 2025).

Использование и адаптация существующих моделей управления рисками к проектам AI является важным условием полноты и последовательности формирования системы риск-менеджмента. Так, для организации управления рисками и распределения ролей и обязанностей может применяться трехуровневая модель защиты, предусматривающая разделение функций управления, контроля и независимого аудита на различных уровнях. Первый уровень защиты связан с предоставлением AI-продуктов

⁴⁰ Например, число зарегистрированных рисков и инцидентов, связанных с применением GenAI, увеличилось с 70 в марте 2023 г. до 650 в сентябре 2024 г. (OECD, 2023).

⁴¹ К числу таких организаций относятся Организация экономического сотрудничества и развития (OECD), Международная организация по стандартизации (ISO) и др. Внедрение подобных принципов должно производиться на адаптивной основе с учетом операционной среды и склонности к риску каждой организации. Кроме того, излишнее нормативно-правовое регулирование AI в нынешних условиях может не столько способствовать снижению рисков, сколько тормозить развитие инноваций и снижать темпы адаптации технологий искусственного интеллекта в финансовой сфере. Последнее имеет непосредственное отношение к вступившему в силу 01.08.2024 г. Закону Евросоюза об искусственном интеллекте (The European Union's Artificial Intelligence Act) (Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 Laying Down Harmonised Rules on Artificial Intelligence and Amending Regulations (EC) (Artificial Intelligence Act). *Official Journal of the European Union*, OJ L. 12.7.2024).

⁴² OECD (2024). Recommendation of the Council on Artificial Intelligence. OECD Legal Instruments. <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449> (accessed on May 10, 2026)

⁴³ В финансовом регулировании ключевая задача объяснимости AI заключается не в достижении полной прозрачности моделей, а в обеспечении достаточной и риск-ориентированной интерпретируемости их решений (Pérez-Cruz et al., 2025).

⁴⁴ Подробнее см.: BIS, 2025.

⁴⁵ Consultative Group on Risk Management (2023). Central Bank Digital Currency (CBDC) information security and operational risks to central banks, November. <https://www.bis.org/publ/othp81.pdf> (accessed on May 10, 2026)

и услуг внутренним и внешним пользователям. На этом уровне осуществляются: определение случаев использования AI, соответствующих стратегии ЦБ или обладающих инновационной ценностью; идентификация и оценка рисков с учетом специфических угроз и уязвимости AI; применение технических мер контроля (тестирование надежности моделей, проверка обучающих данных и результатов, методы выявления предвзятости и др.); настройка механизмов мониторинга для выявления проблем производительности и безопасности; а также обучение персонала, направленное на обновление навыков и компетенций в области искусственного интеллекта и управления рисками.

Второй уровень защиты отвечает за регулирование и надзор за стандартами AI, поддерживая первый уровень и при необходимости пересматривая его оценки с целью повышения эффективности контроля. На данном уровне осуществляются: согласование использования AI с риск-аппетитом центрального банка; разработка методологии управления рисками и координация риск-оценок первого уровня, а также приоритизация рисков на основе профиля рисков AI; обновление существующих политик и разработка новых руководящих принципов для устранения рисков AI; контроль соблюдения нормативных и этических требований, включая анализ правовой базы с учетом специфики финансового сектора и деятельности ЦБ; а также подготовка и распространение учебных материалов по рискам AI для повышения осведомленности персонала.

Третий уровень защиты обеспечивает независимую и объективную оценку адекватности и эффективности управления, а также риск-менеджмента. На этом уровне реализуются следующие функции. Во-первых, осуществляются технические и этические аудиты, в ходе которых оцениваются меры безопасности, а также справедливость, прозрачность и этичность моделей AI. Во-вторых, проводятся проверки и разрабатываются рекомендации по улучшению мер контроля в области AI с учетом быстро развивающихся технологий и связанных с ними рисков (BIS, 2025).

Использование трехуровневой модели защиты обладает рядом преимуществ в управлении рисками AI, поскольку четкое распределение ролей и обязанностей на каждом уровне облегчает интегрированное управление рисками, повышает способность оперативно выявлять и снижать угрозы, а также усиливает устойчивость к потенциальным инцидентам. При этом система управления рисками AI в центральных банках должна основываться на принципах, ключевым из которых является внедрение надежной программы управления ИКТ-рисками в рамках системы управления операционными рисками, что способствует повышению операционной устойчивости. Ключевыми элементами такой системы являются операционно-цифровое управление рисками, обеспечение непрерывности бизнеса, картирование и контроль зависимостей, управление инцидентами, управление рисками информационной безопасности, а также обеспечение кибербезопасности.

В связи с большим объемом данных, используемых моделями AI, и ограниченной прозрачностью ряда алгоритмов, одной из ключевых задач управления рисками является защита информации в части конфиденциальности, целостности, доступности и неприкосновенности частной жизни. Соответственно, ЦБ должны не только предотвращать кибератаки и защищать ИКТ-инфраструктуру, но и обеспечивать способность оперативного восстановления и поддержания критически важных функций в условиях киберинцидентов. С учетом вышеизложенного для минимизации конкретных рисков внедрения AI в центральных банках может быть предложен ряд мер.

Для минимизации операционного риска ЦБ следует определить организационные функции, связанные с внедрением AI, включая формирование культуры управления рисками на всех этапах жизненного цикла систем искусственного интеллекта, а также оценку, мониторинг и реагирование на риски с применением мер по снижению их негативных последствий. Необходимо четко распределить роли и обязанности между владельцами, разработчиками, пользователями и валидаторами моделей AI. Также следует установить требования к AI и сформировать политику его внедрения, определяющую допустимые и запрещенные практики использования искусственного интеллекта⁴⁶.

Для минимизации стратегического риска необходимо привести использование AI в соответствие с нормативными требованиями и интегрировать его в процессы мониторинга и отчетности. Нор-

⁴⁶ Запрещенными практиками могут являться: 1) публикация документов, созданных с использованием AI, без надлежащей проверки человеком, особенно в условиях ограниченного времени; 2) использование моделей и инструментов AI в критически важных процессах до их надлежащей калибровки и валидации; 3) использование конфиденциальной информации для обучения моделей AI без соответствующих механизмов защиты данных; 4) предоставление системам AI возможности автономно модифицировать собственные алгоритмы без контролируемых процедур и надзора (BIS, 2025).

мативные требования должны также учитываться при обработке конфиденциальной информации и персональных данных в качестве входных данных для моделей AI. Для минимизации рисков информационной безопасности, конфиденциальности и кибербезопасности следует классифицировать модели AI в зависимости от уровня чувствительности данных, используемых для их обучения, и определить допустимые типы информации на основе данной классификации. Необходимо ограничить доступ к публичным онлайн-моделям, использовать изолированные среды (сети или «песочницы») для AI-инструментов, а также минимизировать использование конфиденциальных данных при формировании входных запросов. Дополнительно следует учитывать специфические угрозы, включая инъекции подсказок, уязвимости приложений, небезопасный дизайн плагинов и отравление обучающих данных, а также выявлять и устранять уязвимости конкретных LLMs.

В целях минимизации ИКТ-рисков следует анализировать текущие и будущие потребности в ресурсах, включая вычислительную мощность, объемы хранения данных и пропускную способность информационных систем ЦБ. Необходимо обеспечить соответствие существующей инфраструктуры нагрузкам, связанным с использованием AI, а также осуществлять постоянный мониторинг и оптимизацию распределения ресурсов для повышения эффективности и производительности. Для снижения риска третьей стороны и репутационных рисков необходимо изучать и регулярно пересматривать условия предоставления услуг провайдеров во избежание конфликтов интересов и операционных сбоев. Также необходимо привлекать специалистов, обладающих компетенциями в области обработки данных, используемых в системах искусственного интеллекта, особенно в случаях, когда провайдеры хранят и используют данные для обучения моделей AI или передают их третьим лицам. Дополнительно требуется оценивать состав и качество данных, используемых для обучения моделей AI, и применять в деятельности только интерпретируемые результаты их работы.

Несмотря на то, что современные модели AI обучаются на больших объемах данных, полученные результаты следует рассматривать как промежуточные и подлежащие верификации с привлечением профильных экспертов и специализированных оценочных процедур. В качестве основы доверия и подотчетности AI-моделей необходимо сочетание методов объяснимого искусственного интеллекта (ХAI), независимой валидации и человеческого контроля (Pérez-Cruz et al., 2025)⁴⁷. Дополнительно может потребоваться обучение персонала для понимания рисков и ограничений моделей AI, включая генеративные системы, что является необходимым условием их широкого применения в ЦБ. Также важно формировать качественные входные запросы для повышения надежности результатов и обеспечивать маркировку процессов и результатов, в которых использовались AI-системы.

Хотя AI пока не входит в число приоритетных направлений большинства ЦБ по сравнению с цифровой трансформацией и развитием платежных систем⁴⁸, при надлежащем внедрении и управлении рисками в ближайшие годы можно ожидать существенного расширения применения технологий искусственного интеллекта в различных направлениях деятельности центральных банков.

Выводы

На протяжении последних десятилетий технологии искусственного интеллекта существенно эволюционировали от простых вычислительных моделей к современным методам машинного обучения, в том числе GenAI. Ключевым этапом этой эволюции стало развитие архитектуры трансформеров, позволившей существенно повысить качество обработки последовательностей данных и создать LLMs, способные генерировать связный контент и автоматизировать широкий спектр задач.

Развитие технологий искусственного интеллекта открыло новые возможности их применения в деятельности центральных банков. Современными областями применения AI в ЦБ являются экономический анализ, прогнозирование, платежные системы, регулирование, надзор, выявление аномалий и оценка рисков. В целом применение AI повышает эффективность борьбы с отмыванием денег, укрепляет кибербезопасность и поддерживает выполнение регуляторных мандатов в области ценовой и финансовой стабильности.

⁴⁷ Дополнительно могут применяться универсальные или специализированные программно-сетевые решения, в том числе основанные на облачных платформах, предназначенных для разработки доверенных интеллектуальных систем (см.: Турдаков и др., 2022).

⁴⁸ В настоящее время тематика развития AI указывается лишь в 10% стратегических планов ЦБ (см.: Digital transformation and payments are major strategic goals. Central Banking, 2025. <https://www.centralbanking.com/benchmarking/strategic-planning/7972763/digital-transformation-and-payments-are-major-strategic-goals> (accessed on May 10, 2026))

Наиболее активное применение искусственного интеллекта наблюдается в банковском надзоре, экономическом анализе и корректировке протоколов монетарной политики. В совместных проектах с частными финансовыми организациями AI используется для прогнозирования макроэкономики, повышения киберустойчивости и выявления трансграничных финансовых преступлений, где модели машинного обучения превосходят традиционные подходы во многих юрисдикциях.

Одновременно внедрение AI в ЦБ ведет к росту операционного, ИКТ-риска, риска конфиденциальности, риска третьих лиц и репутационного риска, а также к формированию нового типа риска, связанного с моделями AI. В этих условиях внедрение AI должно сопровождаться развитием адаптивных систем управления рисками, обеспечивающих своевременную идентификацию, классификацию и минимизацию потенциальных негативных последствий.

Адаптация существующих моделей риск-менеджмента к проектам AI на основе многоуровневой защиты позволяет разграничить функции управления, контроля и аудита рисков внедрения AI в ЦБ. При надлежащем внедрении и управлении рисками AI в ближайшие годы можно ожидать дальнейшего расширения применения технологий искусственного интеллекта в различных направлениях деятельности ЦБ.

Литература / References

- Болонин А.И., Алиев М.М. (2024). Использование аналитики больших данных и искусственного интеллекта в центральных банках. *Банковские услуги* (5), 12–17. [Bolonin, A. Aliev, M. (2024). The use of big data analytics and artificial intelligence in central banking. *Banking Services* (5), 12–17 (in Russian)]. EDN: QBDTKM
- Гаврилова Э.Н. (2024). Искусственный интеллект в финансовой сфере: эволюция, возможности и перспективы использования. *Вестник Московского университета им. С.Ю. Вумте. Серия 1: Экономика и управление* 50(3), 23–30. [Gavrilova, E. (2024). Artificial intelligence in the financial sphere: Evolution, possibilities and prospects of use. *Economics and management* 50(3), 23–30 (in Russian)]. DOI: 10.21777/2587-554X-2024-3-23-30
- Кочергин Д.А. (2025) Ключевые направления применения искусственного интеллекта в финансовой сфере. *Вестник Института экономики Российской академии наук* (6), 147–169. [Kochergin, D. (2025). Main trends in the use of artificial intelligence in the financial sector. *The Bulletin of the Institute of economics of the Russian academy of science* (6), 147–169 (in Russian)]. DOI: 10.52180/2073-6487_2025_6_147_169
- Турдаков Д.Ю., Аветисян А.И., Архипенко К.В. и др. (2022). Доверенный искусственный интеллект: вызовы и перспективные решения. *Доклады Российской академии наук. Математика, информатика, процессы управления* 508(1), 13–18. [Turdakov, D., Avetisyan, A., Arkhipenko, K. et al. (2022). Trusted artificial intelligence: challenges and promising solutions. *Doklady Mathematics, Computer Science, Control Processes* 508(1), 13–18 (in Russian)]. DOI: 10.31857/S2686954322070207
- Aldasoro, I., Doerr, S., Gambacorta, L. et al. (2024). *Generative artificial intelligence and cybersecurity in central banking*. BIS Papers no. 145, May. <https://www.bis.org/publ/bppdf/bispap145.pdf> (accessed on May 10, 2026)
- Araujo, D., Bruno, G., Marcucci, J., Schmidt, R., Tissot, B. (2022). *Machine learning applications in central banking: An overview in “machine learning in central banking”*. IFC Bulletin no. 57, November. https://www.bis.org/ifc/publ/ifcb57_01_rh.pdf (accessed on May 10, 2026)
- Araujo, D., Schmidt, R., Sirello, O., Tissot, B., Villarreal, R. (2025). *Implementation and governance of artificial intelligence in central banks*. The European Money and Finance Forum. SUERF Policy Brief no. 1185, June. https://www.suerf.org/wp-content/uploads/2025/06/SUERF-Policy-Brief-1185_Araujo_Schmidt_Sirello_Tissot_Villarreal.pdf (accessed on May 10, 2026)
- Araujo, D., Cap, A., Mattei, I., Schmidt, R., Sirello, O., Tissot, B. (2026). *Data science in central banking: Unlocking the potential of data*. IFC Bulletin no. 64, May. https://www.bis.org/ifc/publ/ifcb64_00_rh.pdf (accessed on May 10, 2026)

- Baily, M., Byrne, D., Kane, A., Soto, P. (2025). *Generative AI at the crossroads: Light bulb, dynamo, or microscope?* Finance and Economics Discussion Series no. 2025-053. Federal Reserve Board, Washington, D.C. <https://www.federalreserve.gov/econres/feds/files/2025053pap.pdf> (accessed on May 10, 2026)
- Bank for International Settlements (BIS) (2025). *Governance of AI adoption in central banks*. Consultative Group on Risks Management Paper, 2025, January. <https://www.bis.org/publ/othp90.pdf> (accessed on May 10, 2026)
- Bell, S., Gadanecz, B., Gambacorta, L., Perez-Cruz, F., Shreeti, V. (2025). Artificial intelligence and human capital: Challenges for central banks. *BIS Bulletin* (100), April. <https://www.bis.org/publ/bisbull100.pdf> (accessed on May 10, 2026)
- Boyd, D., Crawford, K. (2012). Critical questions for big data. *Information, Communication & Society* **15**(5), 662–679. DOI: 10.1080/1369118X.2012.678878
- Dauphin, J-F., Dybczak, K., Maneely, M. et al. (2022). *Nowcasting GDP: A scalable approach using DFM, machine learning and novel data, applied to European economies*. IMF Working Paper no. 52, March.
- Dunn, W., Meade, E., Sinha, N., Kabir, R. (2024). *Using generative AI models to understand FOMC monetary policy discussions*. FEDS Notes no. 2024-12-06-1, December.
- Gambacorta, L., Kharroubi, E., Mehrotra, A., Oliviero, T. (2025). *Artificial intelligence and growth in advanced and emerging economies: Short-run impact*. BIS Working Papers no. 1321, December. <https://www.bis.org/publ/work1321.pdf> (accessed on May 10, 2026)
- Garratt, R., Wilkens, P., Shin, H. (2024). Next generation correspondent banking. *BIS Bulletin* (87), May. <https://www.bis.org/publ/bisbull87.pdf> (accessed on May 10, 2026)
- Geiger, F., Kanelis, D., Lieberknecht, P., Sola, D. (2025). *Monetary-Intelligent Language Agent. (MILA)*. Deutsche Bundesbank Technical Paper no. 01/2025. <https://www.bundesbank.de/resource/blob/855186/89fae2a6abdc3ea6de36abe12147269e/472B63F073F071307366337C94F8C870/2025-01-technical-paper-data.pdf> (accessed on May 10, 2026)
- OECD (2023). *Generative artificial intelligence in finance*. OECD Artificial Intelligence Papers no. 9, December. https://www.oecd.org/content/dam/oecd/en/publications/reports/2023/12/generative-artificial-intelligence-in-finance_37bb17c6/ac7149cc-en.pdf (accessed on May 10, 2026)
- OECD (2026). *Supervision of artificial intelligence in finance challenges, policies and practices*. Artificial Intelligence Papers no. 54, January. https://www.oecd.org/content/dam/oecd/en/publications/reports/2026/01/supervision-of-artificial-intelligence-in-finance_1295e5e2/92743dc1-en.pdf (accessed on May 10, 2026)
- Park, T., Shin, H., Williams, H. (2025). *Mapping the space of central bankers' ideas*. BIS Working Papers no. 1299, October. <https://www.bis.org/publ/work1299.pdf> (accessed on May 10, 2026)
- Perez-Cruz, F., Shin, H. (2024). Testing the cognitive limits of large language models. *BIS Bulletin* (83). <https://www.bis.org/publ/bisbull83.pdf> (accessed on May 10, 2026)
- Pérez-Cruz, F., Prenio, J., Restoy, F., Yong, J. (2025). *Managing explanations: How regulators can address AI explainability*. FSA Occasional Paper no. 24, September. <https://www.bis.org/fsi/fsipapers24.pdf> (accessed on May 10, 2026)
- Russell, S., Norvig, P. (2021). *Artificial Intelligence: A Modern Approach*. Pearson.