

Современные источники рисков кибермошенничества: социальная инженерия с применением технологии искусственного интеллекта

Павел Владимирович Ревенков, доктор экономических наук,
профессор кафедры информационной безопасности,
Финансовый университет при Правительстве Российской Федерации,
4-й Вешняковский проезд, д. 4, корп. 2, г. Москва, 109456, Россия.
E-mail: PVRevenkov@fa.ru
<https://orcid.org/0000-0002-0354-0665>

Диана Олеговна Платова, студент 1-го года обучения магистратуры
направления «Управление информационной безопасностью»,
Финансовый университет при Правительстве Российской Федерации,
4-й Вешняковский проезд, д. 4, корп. 2, г. Москва, 109456, Россия.
E-mail: platovadiana31@gmail.com
<https://orcid.org/0009-0001-1277-0171>

Аннотация.

Предмет/тема. В статье рассмотрены наиболее актуальные схемы кибермошенничества с применением методов социальной инженерии и возможностей искусственного интеллекта.

Цели/задачи. Проанализировать причины высокой эффективности мошеннических схем с использованием методов социальной инженерии, включая особенности применения генеративного искусственного интеллекта. Выявить основные уязвимости пользователей, эксплуатируемые злоумышленниками.

Методология. Объектом исследования выступила кредитно-финансовая сфера Российской Федерации. В статье рассматриваются современные тенденции развития генеративного искусственного интеллекта и источники сопутствующих рисков. Основное внимание уделено использованию данной технологии в мошеннических целях. Авторы рассматривают вопросы эволюционного изменения социальной инженерии на основе использования мошенниками возможностей искусственного интеллекта приводят наиболее распространённые схемы мошенничества.

Результаты. Подробно проанализированы схемы мошенничества с применением методов социальной инженерии и возможностей искусственного интеллекта. Представлены новые формы эволюция социальной инженерии, а также правовые и этические аспекты противодействия мошенничеству с использованием технологий искусственного интеллекта.

Научная новизна. В статье выявлены и систематизированы основные способы совершения мошенничества с применением методов социальной инженерии и возможностей искусственного интеллекта. Представлены экономические последствия распространения мошенничества с использованием технологий искусственного интеллекта и даны основные рекомендации пользователям интернета по минимизации рисков, связанных с кибермошенничеством.

Ключевые слова: генеративный искусственный интеллект, социальная инженерия, кибермошенничество

Modern sources of cyber fraud risks: social engineering using artificial intelligence technology

Pavel V. Revenkov, Doctor of Economic Sciences,
Professor of the Department of Information Security,
Financial University under the Government of the Russian Federation,
4th Veshnyakovsky Proezd, Building 2, House 4, Moscow, 109456, Russia.
E-mail: PVRevenkov@fa.ru

<https://orcid.org/0000-0002-0354-0665>

Diana O. Platova, 1st-year master's student in the field of "Information Security Management", Financial University under the Government of the Russian Federation, 4th Veshnyakovsky Proezd, building 4, building 2, Moscow, 109456, Russia.
E-mail: platovadiana31@gmail.com
<https://orcid.org/0009-0001-1277-0171>

Annotation.

Subject/topic. The article examines the most relevant cyber fraud schemes involving the use of social engineering techniques and the capabilities of artificial intelligence.

Goals/objectives. To analyze the reasons for the high effectiveness of fraudulent schemes using social engineering methods, including the specifics of using generative artificial intelligence. To identify the main vulnerabilities of users exploited by attackers.

Methodology. The object of the study is the credit and financial sector of the Russian Federation. The article examines current trends in the development of generative artificial intelligence and the sources of associated risks. The main focus is on the use of this technology for fraudulent purposes. The authors examine the issues of the evolutionary change in social engineering based on the use of artificial intelligence capabilities by fraudsters and present the most common fraud schemes.

Results. The fraud schemes involving the use of social engineering methods and artificial intelligence capabilities have been analyzed in detail. New forms of the evolution of social engineering, as well as the legal and ethical aspects of combating fraud using artificial intelligence technologies, are presented.

Scientific novelty. The article identifies and systematizes the main methods of committing fraud using social engineering techniques and the capabilities of artificial intelligence. It presents the economic consequences of the spread of fraud using artificial intelligence technologies and provides key recommendations for internet users on how to minimize the risks associated with cyber fraud.

Keywords: generative artificial intelligence, social engineering, cyber fraud

Возможности генеративного искусственного интеллекта

Современный мир невозможно представить без применения цифровых технологий. Одним из значимых факторов является активное использование искусственного интеллекта (ИИ). Рынок генеративного ИИ демонстрирует стремительный рост: по итогам 2025 года его объем достиг 58 млрд руб., что более чем в четыре раза превышает показатель 2024 года (13 млрд руб.). Согласно оценкам, к 2030 году ожидается рост до 778 млрд руб. Характерной тенденцией текущего этапа развития становится глубокая интеграция агентных ИИ-систем в инфраструктуру массовых цифровых сервисов. Примером служит внедрение функций саммаризации поисковой выдачи в таких системах, как «Яндекс» и «Google»: встроенные ИИ-агенты (например, «АлисаAI» в экосистеме «Яндекса») анализируют пользовательский запрос и предоставляют краткое содержание наиболее релевантных источников без необходимости перехода по ссылкам.

ИИ оказал революционное влияние на множество сфер – от медицины до развлечений. Однако, чем больше технологий проникает в повседневность, тем сложнее становится различать, где заканчиваются выгоды и начинается риск.

Представим себе, что некто создает злонамеренное программное обеспечение, способное анализировать поведение пользователей. С помощью сложных алгоритмов ИИ эта система выявляет наиболее уязвимые места в безопасности информации. Мошенник, используя такую программу, может с легкостью составить профиль своей потенциальной жертвы. Данные из социальных сетей, такие как личные интересы, местоположение и круг общения, служат основой для создания контекста, который может показаться вполне правдоподобным. Таким образом, технологии не только открывают новые горизонты для искреннего общения, но и подготавливают почву для манипуляций и обмана.

Одним из таких примеров может служить создание фальшивых аккаунтов, которые выглядят абсолютно реалистично. Мошенники могут использовать технологии глубокого обучения для генерации фотографий, видео и текстов, создающих иллюзию общения с настоящим человеком. Это делает предпринятые меры по безопасности малоэффективными.

С внедрением ИИ мошенники получают мощные инструменты, значительно упрощающие создание и распространение ложной информации. Алгоритмы генерации текстов и изображений позволяют получать контент, который сложно отличить от настоящего. Возможно, что сообщение, полученное вами от друга, на самом деле было создано с помощью ИИ, эмулирующего стилистику и язык вашего собеседника. Создание фальшивых профилей стало проще благодаря технологиям глубокого обучения. Генерация изображений лиц, которые фактически не существуют, больше не является чем-то фантастическим. Эти технологии позволяют мошенникам создавать правдоподобные аккаунты в социальных сетях, которые выглядят абсолютно реальными. Всё это формирует новую реальность, в которой идентичность становится всё более размытой. Это угрожает безопасности пользователей и позволяет манипуляторам устанавливать доверительные отношения с жертвой, создавая иллюзии близости и доверия.

Эволюция социальной инженерии

По мере того как технологии ИИ интегрируются в ежедневные задачи миллионов пользователей, наблюдается рост сопутствующих рисков. Увеличение пользовательской базы генеративных моделей закономерно обостряет проблемы конфиденциальности: вопросы сбора, хранения и обработки персональных данных приобретают критическое значение в отсутствие прозрачных механизмов контроля. Одновременно с этим формируется принципиально новый вектор угроз, связанный с применением технологий ИИ в мошеннических целях.

Особую опасность представляет совместное применение возможностей ИИ и социальной инженерии (СИ)¹. СИ, как дисциплина манипуляции, всегда была и остается в движении, адаптируясь к изменениям в обществе и технологиях. В условиях активного использования ИИ эта адаптация становится особенно ощутимой. Новые инструменты, методы обработки данных и алгоритмы машинного обучения открывают перед мошенниками невиданные ранее горизонты возможностей. Существующие техники манипуляции становятся более изощренными и трудными для распознавания.

Одним из наиболее значимых источников сопутствующих рисков можно назвать нарушение конфиденциальности данных пользователей.

ИИ-системы могут хранить всю информацию о каждом пользовательском запросе, формируя из них цифровой профиль человека, дообучать на персональных данных пользователей модели работы нейросетей.

Даже если чат-бот забывает диалоги, персональные данные, полученные в ходе переписки, могут остаться в обучающей выборке, и потом всплыть при последующем обучении модели или генерации ответов другим пользователям, особенно если выборка не была анонимизирована.

Достаточно вспомнить случаи масштабных утечек данных пользователей ChatGPT. В 2023 году из-за технической ошибки часть пользовательской информации стала доступна третьим лицам. Кроме того, отдельные беседы пользователей попадали в результаты индексации поисковых систем, что создавало риск несанкционированного доступа к конфиденциальной информации.

Применение технологий ИИ обеспечивает злоумышленникам еще два существенных преимущества. Во-первых, снижается порог входа в данную сферу за счет автоматизации

¹ Социальная инженерия (social engineering) – это любая атака, которая использует человеческую психологию для воздействия на цель, заставляя ее либо выполнить нужное действие, либо предоставить секретную информацию. Эти атаки играют важную роль в индустрии информационной безопасности и хакерском сообществе, но мы регулярно встречаем примеры подобного поведения в своей повседневной жизни.

рутинных процессов. Во-вторых, мошеннические схемы становятся более масштабируемы: у злоумышленников появляется возможность одновременно воздействовать на несколько жертв без потери убедительности.

Для генерации правдоподобных аудиосообщений злоумышленники могут прибегать к такому инструменту как взлом учетных записей пользователей в мессенджерах. Интерес для них представляют именно голосовые сообщения, накопленные в истории переписки. Жертва переходит по фишинговой ссылке, не осознавая этого, авторизуется на поддельном сайте, и злоумышленники таким образом получают доступ к ее голосовым сообщениям. Сгенерированные на основе собранных аудиозаписей голосовые сообщения с просьбой о срочном переводе денежных средств рассылаются по списку контактов жертвы. При этом жертвы мошенников воспринимают их как подлинные в силу того, что голос является узнаваемым и ему доверяют, что позволяет злоумышленникам успешно обходить традиционные меры защиты (рис. 1).

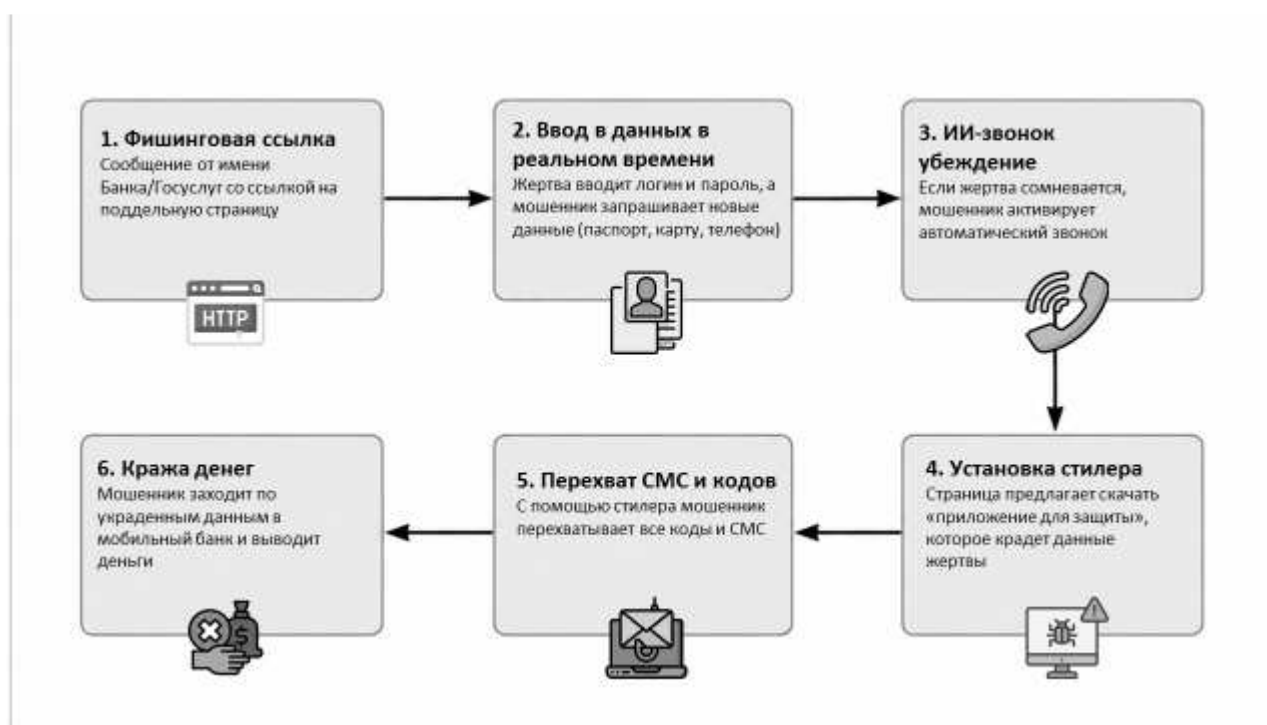


Рисунок 1. Схема мошенничества «Фишинг с использованием искусственного интеллекта»
Источник: составлено авторами

Также нейросети можно обучить манере поведения любого человека на основе его сообщений, публикаций в социальных сетях и других цифровых следов. Подобные технологии могут использоваться мошенниками для создания убедительных поддельных сообщений, голосовых или текстовых обращений от имени знакомых, коллег или родственников жертвы. Это повышает эффективность методов СИ, поскольку злоумышленники способны имитировать характерный стиль общения человека, вызывая больше доверия и снижая вероятность распознавания обмана.

Правовые и этические аспекты противодействия ИИ-мошенничеству

Развитие генеративного ИИ опережает формирование нормативно-правовой базы, что создает существенные пробелы в регулировании. Действующее законодательство в сфере персональных данных ориентировано на классические информационные системы и не в полной мере учитывает специфику нейросетевых моделей, которые могут непреднамеренно «запоминать» и воспроизводить конфиденциальную информацию. Отсутствие четких требований к анонимизации обучающих выборок, срокам хранения пользовательских данных

и обязательному информированию о фактах утечек делает пользователей уязвимыми даже при использовании легальных сервисов.

Кроме того, возникает этическая дилемма: кто несет ответственность за вред, причиненный сгенерированным контентом, – разработчик модели, конечный пользователь или сам алгоритм? Большинство юрисдикций склоняются к тому, что ответственность лежит на лице, использующем ИИ в противоправных целях, однако доказательная база в таких делах крайне сложна. Дипфейки и синтезированные голоса практически не оставляют криминалистических следов, что затрудняет идентификацию злоумышленника.

На международном уровне обсуждается введение обязательной маркировки контента, созданного ИИ, но такие механизмы пока не стали обязательными для массовых платформ. В России Банк России и Роскомнадзор разрабатывают рекомендации по безопасному использованию генеративных моделей, однако они носят скорее рамочный характер. Для эффективного противодействия необходимо внесение изменений в Уголовный кодекс Российской Федерации в части квалификации мошенничества с использованием технологий ИИ как отягчающего обстоятельства, а также создание специализированных подразделений по расследованию киберпреступлений с применением ИИ.

Экономические последствия распространения ИИ-мошенничества

Масштабы финансовых потерь от мошенничества с использованием социальной инженерии и ИИ приобретают макроэкономическое значение. По оценкам Банка России, в 2025 году объем несанкционированных операций без добровольного согласия клиентов вырос на 35% по сравнению с предыдущим годом, причем значительная доля приходится на атаки, использующие синтезированный голос и дипфейки. Эти потери ложатся не только на частных лиц, но и на финансовые институты, вынужденные компенсировать ущерб клиентам в рамках регуляторных требований, что в конечном итоге увеличивает стоимость банковских услуг для всех потребителей.

Бизнес-сектор также становится мишенью: атаки на корпоративные аккаунты в мессенджерах с подменой голоса руководителя позволяют инициировать несанкционированные переводы крупных сумм. В 2025 году зафиксированы случаи, когда такие атаки проводились на компании среднего и малого бизнеса, не имеющих достаточных ресурсов для внедрения многофакторной аутентификации и систем поведенческого анализа.

Кроме прямых финансовых потерь, существует косвенный ущерб: снижение доверия к цифровым каналам связи и электронным платежным инструментам. Это замедляет переход на безналичные расчеты и цифровые сервисы, что противоречит стратегии цифровизации экономики. В долгосрочной перспективе рост ИИ-мошенничества может создать «цифровой разрыв» между поколениями и социальными группами – менее защищенные слои населения будут вынуждены отказываться от современных технологий, что снизит их экономическую активность и доступ к финансовым услугам.

Основные рекомендации пользователям интернета

В сложившихся условиях для снижения факторов риска нарушения кибербезопасности (в том числе СИ с использованием возможностей ИИ) необходимо непрерывно повышать уровень киберграмотности населения. Критический подход к информации, которую мы воспринимаем, а также осведомленность о методах манипуляции становятся неперенными условиями безопасности в сети.

Цифровая грамотность населения стоит в центре борьбы с новыми волнами СИ, но без надлежащего образования и информирования пользователей о рисках и методах защиты многие остаются уязвимыми. Важно научить людей распознавать признаки потенциальной угрозы, развивая критическое мышление и способности к анализу информации. Например, необходимость проверки источников и обращение к факторам, указывающим на мошенничество, – это навыки, которые должны стать частью повседневной практики каждого интернет-пользователя.

Современные тренды в области СИ диктуют необходимость улучшения образовательных программ о кибербезопасности. Постепенно нарастает понимание, что знание методов манипуляции и рисков общения в сети – это не просто привилегия, а необходимость. Мы должны научиться распознавать не только мошенников, но и предупреждать свои живые связи о потенциальных угрозах. Это позволит не только снизить степень уязвимости отдельных пользователей, но и повысить общий уровень кибербезопасности в обществе.

Литература:

1. Ревенков П. В., Бердюгин А. А. Социальная инженерия как источник рисков в условиях дистанционного банковского обслуживания // Национальные интересы: приоритеты и безопасность. – 2017. – Т. 13, № 9. – С. 1747–1760.
2. Банк России. Кибермошенничество: противодействие новым угрозам [Электронный ресурс] // Центральный банк Российской Федерации. – URL: <https://cbr.ru/Content/Document/File/174841/CYBERFRAUDCOUNTERINGNEWTHTREATS.pdf> (дата обращения: 11.05.2026).
3. Кибербезопасность в условиях электронного банкинга. Практическое пособие / под ред. П. В. Ревенкова. – М.: Издательство «Прометей», 2020. – 520 с.
4. Обзор операций, совершенных без добровольного согласия клиентов финансовых организаций [Электронный ресурс] // Центральный банк Российской Федерации. – 2025. –
5. Ашманов, И. С. Цифровая гигиена / И. С. Ашманов, Н. И. Касперская. – Санкт-Петербург : Питер, 2021. – 398 с.
6. Шейнов, В. П. Психология обмана и мошенничества / В. П. Шейнов. – Москва : АСТ, 2007. – 463 с.
7. Рябцев, И. Г. Психология мошенничества и обмана. Как думают и действуют мошенники и как мы можем их распознать / И. Г. Рябцев. – Москва : Эксмо, 2026. – 301 с.
8. Козлов, Н. И. Социальная инженерия. Как нас заставляют делать то, чего мы не хотим / Н. И. Козлов. – Москва : Эксмо, 2024. – 170 с.
9. Грей, Д. Социальная инженерия и этичный хакинг на практике / Д. Грей ; пер. с англ. В. С. Яценкова. – Москва : ДМК Пресс, 2023. – 228 с. : ил.
10. Chatterji A., Cunningham T., Deming D., Hitzig Z., Ong C., Shan C., Wadman K. How People Use ChatGPT [Электронный ресурс]. – OpenAI, 2025. – URL: <https://cdn.openai.com/pdf/a253471f-8260-40c6-a2cc-aa93fe9f142e/economic-research-chatgpt-usage-paper.pdf> (дата обращения: 30.05.2026).

References

1. Revenkov P. V., Berdyugin A. A. Social engineering as a source of risks in the context of remote banking services // National Interests: Priorities and Security. – 2017. – Vol. 13, No. 9. – Pp. 1747–1760.
2. Bank of Russia. Cyber fraud: countering new threats [Electronic resource] // Central Bank of the Russian Federation. – URL: <https://cbr.ru/Content/Document/File/174841/CYBERFRAUDCOUNTERINGNEWTHTREATS.pdf> (accessed: 11.05.2026).
3. Cybersecurity in the context of electronic banking. A practical guide / edited by P. V. Revenkov. – Moscow: Prometheus Publishing House, 2020. – 520 p.
4. Overview of transactions conducted without the voluntary consent of clients of financial organizations [Electronic resource] // Central Bank of the Russian Federation. – 2025. – URL: https://cbr.ru/analytics/ib/operations_survey/2025/ (accessed: 11.05.2026).
5. Ashmanov, I. S. Digital Hygiene / I. S. Ashmanov, N. I. Kasperskaya. – Saint Petersburg: Piter, 2021. – 398 p.

6. Sheynov, V. P. Psychology of Deception and Fraud / V. P. Sheynov. – Moscow: AST, 2007. – 463 p.
7. Ryabtsev, I. G. Psychology of Fraud and Deception. How Scammers Think and Act, and How We Can Recognize Them / I. G. Ryabtsev. – Moscow: Eksmo, 2026. – 301 p.
8. Kozlov, N. I. Social Engineering. How We Are Forced to Do What We Don't Want / N. I. Kozlov. – Moscow: Eksmo, 2024. – 170 p.
9. Gray, D. Social Engineering and Ethical Hacking in Practice / D. Gray; translated from English by V. S. Yatsenkov. – Moscow: DMK Press, 2023. – 228 p.: ill.
10. Chatterji A., Cunningham T., Deming D., Hitzig Z., Ong C., Shan C., Wadman K. How People Use ChatGPT [Electronic resource]. – OpenAI, 2025. – URL: <https://cdn.openai.com/pdf/a253471f-8260-40c6-a2cc-aa93fe9f142e/economic-research-chatgpt-usage-paper.pdf> (accessed: 30.05.2026).