

Осторожно, цифровой обман: как мошенники используют возможности искусственного интеллекта

Еще несколько лет назад искусственный интеллект казался чем-то из будущего. Сегодня он помогает писать тексты, создавать изображения, переводить документы и искать информацию. Но, как это часто бывает с новыми технологиями, их возможности заинтересовали не только разработчиков и бизнес, но и мошенников.

Когда доверие к ChatGPT играет на руку мошенникам

Если раньше злоумышленникам приходилось тратить недели на подготовку убедительных писем, создание фальшивых сайтов или подбор легенд для телефонных звонков, то теперь многие из этих задач можно решить за считанные минуты с помощью нейросетей. В результате схемы обмана становятся все более правдоподобными, а распознать их становится сложнее.

Недавний случай, обнаруженный специалистами по кибербезопасности, показывает, насколько изобретательными становятся преступники. Исследователи компании Push Security обнаружили новую схему **распространения вредоносного программного обеспечения**. Злоумышленники начали использовать функцию общих ссылок ChatGPT, чтобы размещать поддельные страницы с сообщениями о якобы возникших сбоях в работе сервиса. Пользователь видел знакомый домен ChatGPT и потому не подозревал подвоха. Затем ему предлагали скачать «официальное приложение OpenAI» для решения проблемы. На деле под видом программы распространялось вредоносное ПО.

Схема построена на простом психологическом механизме. Большинство людей привыкли считать известные сервисы безопасными. Когда человек видит знакомый логотип или адрес сайта, его бдительность снижается. Этим и воспользовались мошенники.

Подобные атаки **опасны тем**, что они не требуют взлома самого сервиса. Мошенники используют доверие пользователей к популярному бренду и превращают его в инструмент социальной инженерии.

Начальник оказался ненастоящим...

Еще одно направление, которое активно развивается благодаря искусственному интеллекту, создание поддельных голосов или изображений (дипфейк). Сегодня злоумышленникам достаточно нескольких аудио- или видеозаписей из открытых источников, чтобы с помощью искусственного интеллекта **создать цифровую копию** человека. Такие технологии позволяют подделывать голос, внешность и манеру общения, делая мошеннические звонки и видеосообщения максимально убедительными.

Подобные схемы уже стали **причиной многомиллионных убытков** компаний по всему миру. Используя технологии искусственного интеллекта, злоумышленники создают цифровую копию голоса и внешности руководителя, а затем связываются с сотрудниками финансовых подразделений по телефону или через видеосвязь. Под предлогом срочности они просят перевести деньги, оплатить счет или провести другую финансовую операцию. Благодаря высокой реалистичности таких подделок мошеннические звонки и видеосообщения могут не вызывать подозрений даже у опытных специалистов.

Для обычных граждан риски связаны прежде всего с использованием голоса и изображения родственников, друзей или знакомых. Мошенники могут подделать видеозвонок или голосовое сообщение и сообщить о якобы произошедшем ДТП, задержании правоохранительными органами, проблемах со здоровьем или другой чрезвычайной ситуации, требующей срочной финансовой помощи.

Эксперты отмечают, что современные технологии позволяют создавать все более реалистичные аудио- и видеоподделки. В результате распознать фальшивый голос или изображение становится все сложнее. Более того, результаты ряда исследований показывают, что люди нередко ошибаются при попытке отличить сгенерированные материалы от настоящих.

Собеседование, которого никогда не было

Рынок труда тоже оказался под ударом.

В последние годы специалисты по кибербезопасности фиксируют рост мошеннических вакансий, созданных с помощью искусственного интеллекта. Нейросети позволяют быстро генерировать реалистичные описания должностей, сайты компаний и переписку от имени рекрутеров. Соискателю предлагают пройти онлайн-собеседование, заполнить анкеты или установить программное обеспечение для удаленной работы. В результате злоумышленники **получают персональные данные** либо **заражают устройство жертвы** вредоносными программами.

Иногда схема становится еще сложнее. Под видом работодателей преступники проводят целую серию интервью, изучают резюме кандидата и выстраивают доверительные отношения. Только после этого человеку отправляют ссылку на тестовое задание или специальную программу, которая **оказывается вредоносной**.

Парадоксально, но искусственный интеллект используется и, с другой стороны. Некоторые злоумышленники создают полностью вымышленных кандидатов для трудоустройства в компании. Они генерируют фотографии, резюме и даже участвуют в видео собеседованиях с помощью дипфейков. Целью может быть получение доступа к корпоративным системам и данным.

Почему ИИ делает мошенников опаснее

Главное преимущество искусственного интеллекта для преступников – масштабирование. Если раньше мошенник мог подготовить десятки писем в день, то теперь **нейросеть способна создать тысячи персонализированных сообщений за несколько минут**. Если раньше фальшивые тексты содержали множество ошибок, то сегодня они выглядят почти безупречно. Если раньше подделка голоса требовала сложного оборудования, то теперь подобные сервисы становятся доступнее и дешевле.

Особую опасность **представляет скорость**. Если раньше мошенническая кампания могла готовиться неделями, то сейчас злоумышленники способны практически мгновенно реагировать на громкие события. Произошел сбой популярного сервиса, появилась новость о новых выплатах, изменились банковские правила – уже через несколько часов могут появиться десятки фальшивых сайтов, сообщений и рекламных объявлений, созданных с помощью искусственного интеллекта.

Фактически ИИ снижает порог входа в мошенническую деятельность. Для запуска масштабной схемы больше не обязательно обладать глубокими знаниями в области программирования, дизайна или копирайтинга. Многие задачи автоматизируются, а значит, мошеннические операции становятся дешевле, быстрее и доступнее для более широкого круга злоумышленников.

Именно поэтому специалисты по кибербезопасности все чаще говорят о том, что главным объектом атаки становится не компьютер и не смартфон, а человек. Искусственный интеллект помогает мошенникам создавать все более убедительные сценарии, однако окончательное решение; перейти по ссылке, установить программу или перевести деньги – по-прежнему **принимает пользователь**. Поэтому критическое мышление и цифровая грамотность остаются самой надежной защитой даже в эпоху стремительного развития ИИ.

Как безопасно пользоваться сервисами на основе ИИ

Сам по себе искусственный интеллект **не опасен**. Большинство рисков возникает тогда, когда мошенники используют доверие людей к новым технологиям. Поэтому при работе с ИИ-сервисами важно соблюдать несколько простых правил:

- скачивайте приложения только с официальных сайтов разработчиков и из официальных магазинов приложений;
- не переходите по подозрительным ссылкам даже в том случае, если они содержат знакомые названия сервисов;
- не устанавливайте программы, если источник вызывает хотя бы малейшие сомнения;
- критически относитесь к любым сообщениям о срочных проблемах, блокировках или необходимости немедленных действий;
- если вам звонит родственник или руководитель с просьбой срочно перевести деньги, обязательно свяжитесь с ним по другому каналу связи;

- не передавайте личные данные, коды подтверждения и банковскую информацию через чат-ботов и неизвестные сайты;
- внимательно проверяйте вакансии и работодателей, особенно если речь идет об удаленной работе;
- используйте двухфакторную аутентификацию для важных аккаунтов.

Искусственный интеллект уже стал частью нашей повседневной жизни и будет играть **все большую роль в будущем**. Поэтому сегодня важно учиться не только пользоваться новыми технологиями, но и понимать, как ими могут воспользоваться мошенники.

Цифровая грамотность постепенно становится такой же необходимой привычкой, как умение не открывать дверь незнакомцам или не сообщать посторонним данные банковской карты.

Материал подготовлен при помощи www.fingramota.kz