

Почему прогнозы об искусственном интеллекте устаревают быстрее, чем книги о нем



Дата публикации: 17.08.2026

Искусственный интеллект развивается настолько быстро, что книги о нем могут частично устареть еще до выхода из типографии. То, что вчера выглядело как рекламное преувеличение технологических компаний, через несколько месяцев иногда превращается в рабочий инструмент. Именно этот разрыв между сегодняшними возможностями ИИ и его потенциальным будущим оказался в центре спора вокруг новой книги писателя и технологического активиста Кори Доктороу.

Доктороу известен критическим отношением к крупным цифровым платформам и тому, как технологии используются для контроля работников. Один из его наиболее наглядных образов — «обратный кентавр». В обычной модели человек определяет цель, а машина помогает ее достичь. В обратной ситуации алгоритм принимает решения, а человек фактически превращается в исполнителя его инструкций. Примером может быть курьер, которому программа задает маршрут, темп работы и даже допустимые перерывы.

Главную опасность современного ИИ Доктору видит прежде всего не в фантастическом сверхинтеллекте, а в автоматизации труда и стремлении компаний сокращать расходы. По его логике, работодателю выгодно заменить часть сотрудников программным обеспечением или оставить человека контролировать алгоритм, выполняющий основную работу.

Такая критика имеет основания. Искусственный интеллект уже используется для создания текстов, программного кода, изображений, обработки документов и выполнения множества офисных операций. Однако спор начинается вокруг вопроса о том, насколько далеко способна зайти эта автоматизация.

Еще недавно программных ИИ-помощников нередко сравнивали с ненадежным автодополнением: они могли написать небольшой фрагмент кода, но регулярно ошибались и требовали постоянной проверки. Теперь системы способны выполнять более продолжительные задачи, искать ошибки в больших проектах и самостоятельно использовать программные инструменты.

Похожая ситуация складывается с ИИ-агентами. В отличие от обычного чат-бота, отвечающего на конкретный запрос, агенту можно поставить цель и поручить самостоятельно выполнить последовательность действий. Такие системы пока остаются ненадежными, однако сам факт их быстрого развития осложняет прогнозирование.

Главная проблема состоит в том, что нельзя уверенно оценивать будущее технологии только по недостаткам ее нынешнего поколения. Но столь же ошибочно считать, что нынешние темпы прогресса обязательно сохранятся и приведут к появлению искусственного сверхинтеллекта.

Особое внимание исследователей привлекает возможность автоматизации разработки самого ИИ. Если будущие модели смогут существенно помогать ученым и инженерам создавать более совершенные системы, развитие технологии потенциально может ускориться. Улучшенный ИИ будет участвовать в создании следующего поколения моделей, а те — помогать разрабатывать еще более мощные системы.

Именно поэтому дискуссия постепенно выходит за пределы вопроса о том, сколько профессий сможет автоматизировать искусственный интеллект. Появляется другой вопрос: как контролировать разработку систем, которые становятся автономнее и способны выполнять все более сложные интеллектуальные задачи.

Одним из возможных подходов становится регулирование не только применения ИИ, но и процесса создания наиболее мощных моделей. Обсуждаются раскрытие целей их обучения, независимое тестирование

возможностей, контроль вычислительных ресурсов и наблюдение за использованием самых производительных специализированных чипов.

Важна и прозрачность разработчиков. Модель, доступная обычным пользователям, не обязательно является самой мощной системой внутри лаборатории. Если экспериментальные разработки заметно превосходят публичные версии, государствам и независимым экспертам трудно оценивать реальную скорость технологического прогресса.

Проблема дополнительно осложняется международной конкуренцией. Если одна страна ограничит применение или разработку ИИ, это не означает, что остальные государства поступят так же. Поэтому меры регулирования должны учитывать глобальный характер индустрии и огромную роль вычислительной инфраструктуры.

Для оценки будущего все чаще используются сценарные прогнозы. Один из известных примеров — проект AI 2027, авторы которого попытались описать возможное развитие искусственного интеллекта по конкретным этапам: рост вычислительных мощностей, развитие ИИ-агентов, конкуренцию лабораторий и появление систем, способных активнее участвовать в исследованиях.

Ценность подобных сценариев заключается не в том, что они гарантированно сбудутся. Важно другое: прогноз формулируется достаточно конкретно, чтобы спустя некоторое время можно было проверить, где его авторы оказались правы, а где ошиблись.

При этом социальные проблемы сегодняшнего ИИ никуда не исчезают. Даже если сверхинтеллект никогда не будет создан, существующие алгоритмы уже способны изменить рынок труда, творчество, образование, информационную среду и способы управления сотрудниками.

Поэтому спор «ИИ — это революция» против «ИИ — это только хайп» постепенно становится слишком примитивным. Реальность может оказаться сложнее обоих сценариев. Искусственный интеллект способен одновременно оставаться несовершенным, совершать серьезные ошибки и при этом достаточно быстро расширять круг задач, где его использование экономически выгодно.

Главная опасность в такой ситуации — чрезмерная уверенность. Скептики рискуют недооценить реальный технологический прогресс, а сторонники ИИ — принять эффектные демонстрации за доказательство неизбежного появления сверхинтеллекта. Обществу придется учитывать сразу несколько возможных вариантов будущего и создавать правила до того, как последствия новой технологии станут полностью очевидными.