

УДК 004.8:81'25

**КАК НЕЙРОСЕТИ МЕНЯЮТ ЯЗЫКОВУЮ ПРАКТИКУ: ПЕРЕВОД,
ОБУЧЕНИЕ И НЕЗАМЕТНЫЕ ПОСЛЕДСТВИЯ**

Г.В. Литвинова, Д.В. Джигова

*Крымский федеральный университет имени В. И. Вернадского
Симферополь, Российская Федерация, dasha_dzhigova@mail.ru*

Рассматриваются ключевые технологические и прикладные сдвиги, вызванные внедрением нейросетевых языковых моделей в переводческое дело и преподавание иностранных языков. Прослежена эволюция автоматической обработки текста – от

первых систем на жёстких правилах до больших генеративных архитектур. Основное внимание уделено неочевидным ограничениям, заложенным в вероятностной логике моделей, включая феномен ложных утверждений и когнитивных смещений.

Ключевые слова: генеративные нейросети; машинный перевод; языковые галлюцинации; лингводидактика; чат-боты; приватность данных; цифровая этика.

Когда я впервые показала своим друзьям студентам-переводчикам возможности тогда ещё нового чат-бота, старый корпус Крымского федерального наполнился смешанным гулом удивления и тревоги: машина перевела сложный публицистический отрывок быстрее, чем они успели дочитать оригинал. Уже через полгода те же студенты приносили на пары примеры, где нейросеть «изобретала» биографии учёных, перевирала даты и путала географию полуострова. Именно этот разрыв – между впечатляющей скоростью и фундаментальной ненадёжностью – стал для меня отправной точкой для систематического разбора того, как большие языковые модели трансформируют нашу профессиональную среду.

Чтобы понять нынешнее положение дел, стоит коротко восстановить логику развития автоматического перевода. Первые системы, появившиеся ещё во второй половине XX века, работали как старательные, но негибкие лаборанты: им давали грамматические схемы и двуязычные словари, и они собирали перевод поблочно [1, с. 51]. Результат был предсказуемым – корявый синтаксис и полная потеря контекста. Затем пришла эпоха статистических методов: программа уже не пыталась осмыслить фразу, а вычисляла наиболее вероятный перевод на основе параллельных корпусов. Это резко повысило гладкость текста, однако связность за пределами трёх-четырёх слов оставалась слабым местом. Настоящий прорыв обеспечила архитектура Transformer с механизмом внимания, которая дала алгоритмам возможность удерживать семантические связи внутри длинных предложений и даже между соседними абзацами [3, с. 205]. На этом фундаменте выросли и специализированные переводные движки типа DeepL, и генеративные гиганты вроде GPT.

Сегодня мы имеем две пересекающиеся, но не идентичные линейки языкового искусственного интеллекта. Узкопрофильные переводческие системы отлично работают с технической документацией – инструкциями, стандартами, отчётами, где важна терминологическая точность. Универсальные же модели берут шире: они способны не только переводить, но и суммаризировать, перефразировать, стилизовать. По данным недавнего сравнительного исследования на материале уйгурской авторской прозы, при тонкой настройке промптов генеративная модель заметно точнее передавала культурный колорит и эмоциональные оттенки, чем стандартный нейросетевой переводчик, который выдавал грамматически безукоризненный, но стерильный вариант [8, с. 5]. Платой за гибкость стала склонность придумывать отсутствующие в оригинале детали: в нескольких случаях модель вставляла в перевод названия несуществующих книг и имена вымышленных исторических персонажей.

В этом месте мы подходим к центральной уязвимости современных языковых моделей – так называемым галлюцинациям. Сам термин, на мой взгляд, неудачен, потому что предполагает искажение воспринимаемой реальности, тогда как нейросеть вообще никакой реальности не воспринимает. Она предсказывает очередной токен, опираясь на распределение вероятностей, выученное на этапе обучения. Математики, анализировавшие эту особенность в 2025 году,

сформулировали жёсткий вывод: генерация ложных утверждений – не ошибка программирования, а неустранимое следствие вероятностной природы предсказания текста [10, с. 15]. Представьте себе: модель тысячу раз видела в обучающем корпусе фразу «Симферополь – город в...» и почти безошибочно продолжит её правильно. Но стоит запросить малораспространённый факт, как накопление крошечных неопределённостей на каждом шаге порождения приводит к синтезу красивого, гладкого и абсолютно ложного ответа. Ситуация усугубляется тем, что в процессе обучения любой конкретный ответ подкрепляется, а осторожное «я не знаю» расценивается как провал. В итоге системе оказывается выгоднее солгать, чем признаться в нехватке данных.

Практические последствия этой особенности я наблюдаю постоянно. На семинарах по литературному переводу я регулярно даю задание: взять краткий пересказ научной статьи, сгенерированный чат-ботом, и сверить с первоисточником. Почти в каждой студенческой работе обнаруживаются вымышленные цифры, неверно приписанные авторства и даже целые фальсифицированные выводы. Хуже того, при попытке исправить ошибку путём уточняющего запроса модель демонстрирует так называемую «ложно-коррекционную петлю»: она многословно извиняется, заверяет, что теперь учла все замечания, и немедленно продуцирует новую, столь же далёкую от истины версию [9, с. 5]. Постредактирование машинного текста в таких условиях превращается в подобие археологической реконструкции: редактор вынужден не просто шлифовать стиль, а отсеивать фантомные пласты информации.

Было бы ошибкой, однако, говорить о языковом ИИ исключительно в критическом ключе. В сфере преподавания, которой я занимаюсь большую часть времени, открылись возможности, о которых ещё пять лет назад приходилось только мечтать. У нас в Крыму, в многоязычной аудитории, где соседствуют носители русского, крымскотатарского и украинского языков, чат-боты-репетиторы стали незаменимым подспорьем. Я не раз замечала, что студенты, которые на очных парах стесняются говорить из-за акцента или боязни ошибиться, с удовольствием общаются с безличным собеседником, делают по двадцать попыток, переспрашивают и в конечном счёте значительно быстрее преодолевают речевой барьер. Свежий метаанализ более сотни эмпирических исследований за 2024–2025 годы подтверждает: регулярная практика с ИИ-тьютором даёт статистически значимый прирост коммуникативной уверенности, причём главным фактором респонденты называют именно снижение тревожности [6, с. 251]. Машина не оценивает, не сравнивает и не устаёт – и это свойство в определённых учебных ситуациях оказывается важнее дидактической проработки.

Вместе с тем я всё чаще выступаю за осторожный оптимизм. Во-первых, подавляющая часть исследований проведена на глобальном англоязычном материале; для миноритарных языков, включая крымскотатарский, сопоставимых данных практически нет. Во-вторых, мы пока крайне мало знаем о том, как длительное общение с бесстрастным алгоритмом влияет на способность считывать живую интонацию, паузу, невербалику собеседника. Сужу по собственным наблюдениям: школьники и студенты, чрезмерно полагающиеся на чат-ботов, порой начинают терять навык спонтанной устной речи – они привыкают, что реплику можно обдумывать минуту и переписывать трижды, а в реальном диалоге такая роскошь отсутствует. Поэтому я убеждена: ключевой

метакомпетенцией сегодняшнего дня становится то, что «алгоритмической грамотностью» – умение не просто пользоваться нейросетевыми инструментами, но критически оценивать полученный результат, перепроверять факты и распознавать синтетический текст [5, с. 45].

Этическая сторона вопроса требует отдельного разговора. Весной 2025 года вскрылся факт несанкционированного использования персональных сообщений для дообучения одной крупной модели – согласия пользователей, естественно, никто не давал [11, с. 4]. Скандал поднял волну обсуждений, но, положив руку на сердце, проблема гораздо масштабнее громких кейсов. Мы не знаем точного состава обучающих корпусов популярных моделей – разработчики ссылаются на коммерческую тайну, – а значит, любой текст, загружаемый в облачный сервис для перевода или редактирования, потенциально может пополнить чей-то тренировочный датасет. В условиях, когда ИИ-инструменты плотно встроены в рабочий процесс и уйти от них практически невозможно, это создаёт серую зону коллективной уязвимости.

Не менее серьёзный вызов – постепенная инфлюация машинного контента в информационное поле. Группа Алессы провела показательный эксперимент: нейросети поручили суммаризировать новостные заметки и замеры эмоциональный сдвиг получившихся аннотаций [2, с. 7]. В 26% случаев тональность выходного текста заметно отклонялась от оригинала, причём чаще всего сложный конфликт упрощался до бинарного противостояния. Ещё любопытнее оказался второй замер: когда тем же испытуемым показывали сгенерированные моделью отзывы на товары, они принимали решение о покупке на треть чаще, чем после знакомства с подлинными покупательскими комментариями. Машинный текст, лишённый человеческих запинок, повторов и грамматических огрехов, парадоксально воспринимается как более убедительный. В мире, где контентные ленты переполнены, это свойство открывает широкие возможности для манипуляции – от потребительской до политической.

Положение усложняется тем, что мировое сообщество пока не выработало единых регуляторных подходов. Сравнительный обзор внутренних этических политик OpenAI, Meta, Google и Microsoft, проведённый Миллером и Сингхом, рисует картину поразительного разноречия: от многоступенчатой фильтрации обучающей выборки до почти полного отказа от модерации под лозунгом «ответственность лежит на пользователе» [12, с. 6]. Отсутствие общих стандартов означает, что при любом инциденте – диффамация, утечка данных, разжигание розни сгенерированным текстом – виновного придётся искать в лабиринте технических и юридических прослоек, что практически гарантирует безнаказанность.

Подводя черту, выделю несколько позиций, которые видятся мне принципиально важными. Первое: граница между специализированными переводчиками и универсальными генераторами текста размывается, и мы движемся к гибридным схемам, в которых черновой перевод выполняет одна система, а финальную отделку – другая. Второе: борьба с галлюцинациями не может быть чисто инженерной задачей, потому что их корень – в математической природе вероятностного предсказания; нужны внешние контуры факт-чекинга и, возможно, изменение самих метрик, по которым оценивается качество модели (поощрение честного отказа от ответа). Третье: в языковом образовании правильная стратегия – не вытеснение преподавателя, а формирование у студента

привычки критического переосмысления машинного результата. Четвёртое: этические риски – приватность, предвзятость, незаметная манипуляция – перешли из теоретической в практическую плоскость, и без скоординированных международных регламентов мы рискуем оказаться в информационной среде, где правду от добротной имитации отличить будет уже почти невозможно.

Если позволить себе заглянуть чуть дальше сегодняшней повестки, главный вопрос, который меня занимает, звучит так: сможем ли мы выработать коллективный иммунитет к синтетическому тексту? И не станет ли это стимулом к тому, чтобы заново осмыслить ценность живого авторства – не как источника однозначной истины, а как носителя уникального человеческого опыта, который никакая нейросеть принципиально не воспроизведёт.

БИБЛИОГРАФИЧЕСКИЕ ССЫЛКИ

1. Алексеева, М. В. От формальных правил к самообучающимся системам: эволюция машинного перевода / М. В. Алексеева // Информационные технологии в гуманитарных науках. – 2024. – № 2. – С. 48–62.
2. Alessa, A. Quantifying cognitive distortion in neural language outputs / A. Alessa, J. T. Brown, W. Chen // Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing. – Singapore, 2025. – P. 1–15.
3. Барановский, П. Д. Архитектура «Трансформер» как лингвистическая революция / П. Д. Барановский, Ю. В. Гнездилова // Вопросы компьютерной лингвистики. – 2025. – № 4. – С. 200–217.
4. Петрова, И. В. Постредактирование нейросетевого перевода: трудозатраты и иллюзия точности / И. В. Петрова // Переводческая практика: наука и искусство. – 2024. – Вып. 18. – С. 81–95.
5. Королёва, Е. Е. Алгоритмическая грамотность как компонент языкового образования / Е. Е. Королёва, Т. А. Смирнов // Вестник Московского университета. Серия 20: Педагогическое образование. – 2025. – № 3. – С. 42–57.
6. A systematic review of AI-driven tools in second language acquisition / N. Lopez, P. Ramírez, S. Chen [et al.] // Computer Assisted Language Learning. – 2025. – Vol. 38, № 2. – P. 245–269.
7. Яковенко, А. С. ChatGPT и DeepL в зеркале литературного перевода: сравнительный анализ / А. С. Яковенко, Л. В. Данилова // Современные парадигмы перевода: сборник научных трудов. – Екатеринбург: Ажур, 2025. – С. 708–715.
8. Wang, Q. Neural rendering of Uighur literature: a comparative output study / Q. Wang // PLoS ONE. – 2025. – Vol. 20, № 10. – e0335261.
9. Structural roots of hallucination in large language models: a model-internal analysis / R. Zimmermann, S. Patel, L. Müller // arXiv preprint arXiv:2511.04567. – 2025. – 23 p.
10. When models guess: mathematical inevitability of AI confabulation / S. Kent, E. Alvarez, T. Fontaine // Nature Computational Science. – 2025. – Vol. 5. – P. 10–19.
11. Проблемы приватности при обучении генеративных нейросетей: аналитический отчёт / Аналитический центр по правам человека. – Минск, 2025. – 28 с.
12. Ethical governance in four tech giants: a comparative study / J. T. Miller, R. A. Singh // AI & Society. – 2025. – Vol. 40. – P. 1–19.

HOW NEURAL NETWORKS ARE RESHAPING LANGUAGE PRACTICE: TRANSLATION, TEACHING AND INVISIBLE CONSEQUENCES

D. V. Dzhigova

*V. I. Vernadsky Crimean Federal University
Simferopol, Russian Federation, dasha_dzhigova@mail.ru*

The paper examines the key technological and applied shifts triggered by the introduction of neural language models in translation and foreign language teaching. The evolution of automatic text processing is traced, from early rule-based engines to large generative architectures. Special attention is given to the inherent limitations of probabilistic generation, including confabulation and cognitive biases. Drawing on classroom experience with philology students in Crimea, the benefits and risks of AI chatbot tutors are discussed. Proposals are formulated for fostering critical digital literacy and ethical standards in language AI environments.

Keywords: generative neural networks; machine translation; AI hallucinations; language teaching; chatbots; data privacy; digital ethics.