

История международных отношений и внешней политики

УДК 004.8:327(73)

DOI: 10.28995/2073-6339-2025-4-12-27

Применение искусственного интеллекта в анализе внешней политики США: обзор современных практик

Наталья А. Цветкова

*Институт США и Канады РАН им. Г.А. Арбатова,
Москва, Россия, n.tsvetkova@iskran.ru*

Аннотация. Статья посвящена анализу направлений использования искусственного интеллекта (ИИ) в исследованиях внешней политики и международных отношений. Автор рассматривает четыре ключевые области применения ИИ: прогнозирование международных событий, моделирование дипломатических переговоров, оценка рисков автоматизированных решений в ядерной сфере и анализ процесса принятия внешнеполитических решений. В центре внимания – сравнительная оценка моделей, разработанных в США, Китае и России, и возможности локальных ИИ как эффективного инструмента экспертного анализа. Особое внимание уделяется методам бенчмаркинга и роли эксперта в интерпретации результатов.

Ключевые слова: внешняя политика, ИИ, дипломатия, прогнозирование, принятие решений, международные отношения, США, стратегическая стабильность

Для цитирования: Цветкова Н.А. Применение искусственного интеллекта в анализе внешней политики США: обзор современных практик // Вестник РГГУ. Серия «Политология. История. Международные отношения». 2025. № 4. С. 12–27. DOI: 10.28995/2073-6339-2025-4-12-27

AI-assisted methods in US foreign policy analysis: A review of current practices

Natalia A. Tsvetkova

*Georgy Arbatov Institute for US and Canada Studies,
Russian Academy of Sciences, Moscow, Russia,
n.tsvetkova@iskran.ru*

Abstract. This article examines the application of artificial intelligence (AI) in the analysis of foreign policy and international relations, with a particular emphasis on the United States. It identifies four principal fields of AI application: forecasting international events, modeling diplomatic negotiations, evaluating the risks of automated decision-making in the nuclear domain, and analyzing foreign policy decision-making processes. The article offers a comparative assessment of AI systems developed in the United States, China, and Russia, underscoring the growing relevance of specialized and domain-specific AI models as tools for expert analysis. Particular attention is devoted to benchmarking methodologies for evaluating generative models and to the critical role of human expertise in interpreting AI-generated outputs.

Keywords: foreign policy, AI, diplomacy, forecasting, decision-making, international relations, US, strategic stability

For citation: Tsvetkova, N.A. (2025), "AI-assisted methods in US foreign policy analysis: A review of current practices", *RSUH/RGGU Bulletin. "Political Science. History. International Relations" Series*, no. 4, pp. 12–27, DOI: 10.28995/2073-6339-2025-4-12-27

Введение

Искусственный интеллект (ИИ) стал неотъемлемым элементом аналитической инфраструктуры в международных отношениях и внешней политике. За последние пять лет наблюдается рост числа научных и прикладных проектов, направленных на использование ИИ для поддержки принятия решений, оценки рисков, моделирования дипломатических сценариев и реконструкции логики внешнеполитического поведения. Понимание того, как именно используются ИИ-модели в этих областях, становится критически важным как для исследователей, так и для практиков международной политики.

ИИ-модель – это обученный алгоритм, способный генерировать текст, классифицировать информацию, выявлять закономерности и строить прогнозы на основе больших массивов данных.

Большинство используемых сегодня систем относятся к категории больших языковых моделей (Large Language Models, LLMs), таких как американские GPT-4o, Llama или китайская Qwen.

Важно различать публичные генеративные чаты, ориентированные на широкий круг задач, и локальные ИИ-системы, разработанные под конкретную тематику – например, внешнеполитический анализ, безопасность или кризисное реагирование. Локальные ИИ чаще всего настраиваются с использованием нужных для экспертов корпусов документов, что повышает точность аналитических выводов и снижает риск ошибок генерации. Для понимания сути работы ИИ необходимо учитывать так называемое обучение ИИ – это процесс обработки больших объемов текстовой и числовой информации, в результате чего модель формирует вероятностные представления о связях между понятиями, событиями и действиями. Алгоритмы и обучающие корпуса документов, использованные в США, Китае и России, существенно различаются, что влияет на характер ответов и интерпретаций, которые генерируют эти модели. В США ИИ применяется преимущественно для поддержки разведки, моделирования рисков и стратегического анализа (проекты StateChat, CamoGPT, Donovan), в Китае – для выработки внешнеполитических решений, согласованных с национальными интересами (Qwen, Doubao, Ernie Bot), а в России ИИ-системы развиваются в рамках оборонных и академических проектов с ограниченным публичным доступом.

Цель данного исследования – выявить существующие подходы к применению искусственного интеллекта в области внешней политики и международных отношений в США. Мы покажем ключевые направления аналитического использования ИИ, а также оценим специфику алгоритмических решений. Особое внимание уделяется возможностям и ограничениям ИИ как аналитического инструмента, влияющего на процессы прогнозирования, принятия решений и моделирования дипломатических сценариев в условиях растущей международной нестабильности.

В настоящей статье выделены ключевые направления исследований в области международных отношений, в которых ИИ находит эффективное применение: прогнозирование международных событий; моделирование дипломатических переговоров; риски автоматизированных решений в ядерной области; процесс принятия внешнеполитических решений. Указанные направления формируют структуру основной части статьи.

Направления исследований с использованием ИИ во внешней политике и международных отношениях

Прогнозирование международных событий. ИИ активно используется для прогнозирования возможных международных событий: вооруженных конфликтов, миграционных кризисов, политической нестабильности. Такие разработки сочетают исторические базы данных и онлайн-источники (новости, спутниковые снимки, социальные сети) с методами прогнозирования.

В первую очередь речь идет о предсказании вероятности вооруженных столкновений, политической нестабильности, гуманитарных кризисов и других форм эскалации. ИИ-платформы анализируют большие массивы данных и выявляют признаки риска. Несмотря на некоторые успехи в прогнозировании общих тенденций, модели нуждаются в значительно большем объеме данных для предсказания конкретных событий. Особенно перспективным признано использование специализированных, локальных ИИ, ориентированных на геополитику конкретных стран (например, китайские модели Qwen и Doubao), которые формируют ответы в соответствии с постулируемыми правительственными ведомствами государственными интересами. Ведется работа по использованию ИИ в моделировании развития ситуации и конечных результатов на основе как исторических аналогий, так и текущих индикаторов. Такие инструменты потенциально позволяют формировать альтернативные сценарии и использовать их в разработке внешнеполитических решений.

Количество открытых публикаций по данной тематике остается ограниченным [Цветкова 2020]. Большинство проектов и разработок осуществляются в закрытом режиме и предназначены для внутреннего использования государственными структурами. В частности ИИ-прогнозирование активно тестируется в Пентагоне и используется в качестве дополнительного инструмента при сценарном планировании. Еврокомиссия заказала исследование по прогнозированию миграционных потоков с помощью ИИ для создания функции «раннего предупреждения» по возможному перемещению населения.

Попытка использовать ИИ для открытых прогнозов в мировой политике предпринята экспертами Калифорнийского университета в Лос-Анджелесе¹. Они разработали базу данных исторических событий и запустили сбор информации для предсказания геополитических событий.

¹ MIRAI Multilingual Open Language Models for Science and Security. 2024. URL: <https://mirai-llm.github.io> (дата обращения 04.03.2025).

литических изменений. ИИ имитирует условия реального мира и обрабатывает современные ситуации. Система продемонстрировала способность к прогнозированию, но при попытке предсказать конкретные типы взаимодействий между странами возникает необходимость в получении дополнительного объема данных. Система слабо справляется с прогнозированием уникальных событий, таких как начало войны или смена власти. Модель не интегрирует данные в реальном времени и не учитывает скрытые политические мотивы и эмоциональные факторы, что ограничивает ее применимость в оперативной аналитике. Наш анализ показывает, что локальный или специализированный ИИ с ограниченными задачами эффективнее решает конкретные вопросы. Однако наличие широкого корпуса информации и его преобразование в датасет является важнейшим этапом насыщения модели для составления аналитических прогнозов.

В силу появления десятков моделей, которые могут использоваться в науке и экспертизе, набирает обороты разработка методик тестирования ИИ на предмет способности рекомендовать экспертам и политикам подходящие внешнеполитические решения. Исследователи предлагают так называемый бенчмаркинг внешнеполитических решений – комплексную оценку того, как разные модели ИИ (американские, китайские или российские) обрабатывают сложные сценарии международной жизни и конфликтов. *Бенчмаркинг* (benchmarking) – это инструмент оценки, позволяющий определить, насколько эффективно модели искусственного интеллекта справляются с задачами, связанными с принятием решений во внешней политике. Оценка состоит из набора стандартных сценариев, на основе которых эксперт сопоставляет ответы различных ИИ-моделей. Цель – выявить отклонения в характере рекомендаций: какие модели предлагают агрессивные, эскалационные действия, а какие – напротив, склонны к дипломатическим или примирительным стратегиям. Бенчмаркинг позволяет оценить предвзятость моделей в критически важных ситуациях и выявить потенциальные риски их применения в стратегическом анализе.

Многие модели склонны предлагать агрессивные меры в кризисных ситуациях. Широко используемые модели ИИ выдают рекомендации по эскалации в кризисных сценариях (рис. 1), причем наибольшую агрессивность показывают широко известные американские модели, такие как “Llama”, и китайская – “Qwen” [Jensen, Reynolds, Atalan 2024]. Это еще раз доказывает, что только эксперт компенсирует ограничения ИИ: даже самая продвинутая модель не может полноценно заменить человеческое стратегическое мышление.

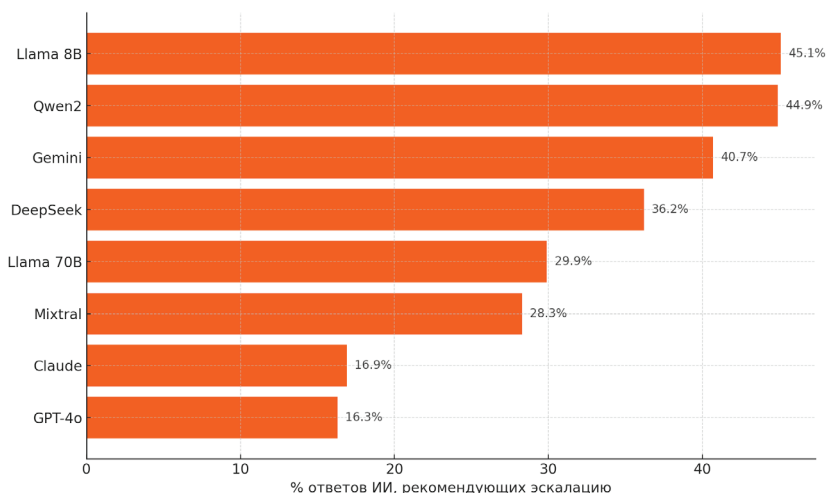


Рис. 1. Рекомендации ИИ по эскалации в кризисных ситуациях

Многие модели ИИ интегрируются в приложения национальной безопасности. В США подобные приложения тестируются в рамках государственных и оборонных структур. В частности, разработаны специализированные ИИ-модели StateChat для анализа внешнеполитических сценариев в Государственном департаменте, а также NIPRGPT и СамоGPT, – в Министерстве обороны для поддержки оперативного анализа и стратегического планирования. В России аналогичные задачи решаются преимущественно в рамках закрытых оборонных и академических программ. Среди них – элементы интеллектуальной обработки данных, разрабатываемые Институтом системного программирования РАН и в рамках оборонных проектов². В Китае активно внедряются национально-ориентированные большие языковые модели, которые применяются в области анализа международных отношений, геополитических рисков и информационного реагирования. Алгоритмические архитектуры, языковые корпуса и источники данных, используемые для обучения этих моделей, существенно различаются. Это определяет вариативность их ответов, интерпретаций и рекомендаций по внешнеполитическим вопросам, что важно при сравнительном анализе результатов моделирования.

² Институт системного программирования им. В.П. Иванникова РАН. URL: <https://www.ispras.ru> (дата обращения 04.03.2025).

Бенчмаркинг показывает разброс вариантов предлагаемых решений – от эскалации до сотрудничества, что и называется предвзятостью в критически важных областях принятия решений во внешней политике. Именно здесь важнейшую роль играет эксперт. Ни одна модель полностью не отражает реальность – даже самый продвинутый ИИ остается приближением, поэтому эксперт выявляет и смягчает ошибки, которые могут исказить стратегический анализ и ввести в заблуждение лидеров национальной безопасности, работающих вместе с ИИ-агентами. При тщательном обучении и подходе к тонкой настройке эти предвзятости могут быть снижены или скорректированы в нужную сторону.

Модели ИИ искажают политический анализ в сторону более агрессивных ответов, что увеличивает риск просчетов в условиях высокой геополитической напряженности. Без постоянной оценки и доработки ИИ-агенты могут усилить склонность к эскалации в сценариях, где сдержанность и стратегическая неопределенность часто являются предпочтительными вариантами³.

Моделирование дипломатических переговоров. Последние исследования показывают, что ИИ может эффективно моделировать дипломатические переговоры, особенно если используется метод RAG – когда модель обращается к проверенной базе документов. Такой подход помогает снизить риск вымышленных ответов и делает сценарии более достоверными. Этот подход подкрепляется анализом баз соглашений и политических текстов: на их основе ИИ формирует возможные сценарии урегулирования, выделяя устойчивые элементы соглашений, ключевые факторы конфликтов и возможные препятствия к переговорам. Использование таких методов, как *chain-of-thought prompting* (поэтапное рассуждение) и *few-shot learning* (обучение по нескольким примерам), помогает выстраивать логику рассуждений, близкую к человеческому анализу, и тем самым повышает реалистичность и аналитическую ценность выводов модели. Результаты показали, что ИИ способен воспроизводить структуры политических решений, предлагать реалистичные и альтернативные варианты соглашений, но не может учесть нестабильность политической воли и непредсказуемые факторы человеческого поведения. Поэтому модели рассматриваются как вспомогательные аналитические инструменты.

³ Critical foreign policy decisions Benchmark // Futures Lab. Center for strategic and international studies. 2025. URL: <https://www.csis.org/programs/futures-lab/projects/critical-foreign-policy-decisions-benchmark> (дата обращения 04.03.2025).

Эксперимент по моделированию условий мирного урегулирования конфликта между Россией и Украиной с помощью ИИ показал несколько моментов [Reynolds, Jensen 2024]. Наличие выборки релевантных документов существенно повышает точность и достоверность ответов, снижая характерную проблему «галлюцинаций» – создания недостоверной или вымышленной информации. Модель таким образом получает возможность опираться на проверенные источники при формировании вывода. Методы *few-shot learning* и *in-context learning*, предполагающие обучение модели с использованием ограниченного количества примеров внутри запроса, позволяет адаптировать модель к специфике задачи. Техника *chain-of-thought prompting* – структурированной постановки задач, позволяет модели держаться логики аргументации и имитировать аналитическое мышление эксперта.

Для обеспечения глубины анализа использовалась тщательно подобранная источниковая база. Основу первого корпуса составили 333 полных текста мирных соглашений, заключенных в период с 1989 по 2018 г. Второй корпус включал 79 аналитических и журналистских материалов, отобранных экспертами, посвященных различным сценариям возможного урегулирования российско-украинского конфликта. Эти данные были загружены в платформу Donovan (<https://scale.com/donovan>) – оболочку, позволяющую настраивать механизмы анализа информации.

В результате моделирования было получено несколько сценариев мирного соглашения, различающихся по политической реалистичности. Первый сценарий включал прекращение огня и вывод войск, но не содержал положений о гарантиях безопасности, статусе спорных территорий или международном мониторинге. Второй сценарий добавлял буферную зону, пакт о ненападении между сторонами и отказ от членства в НАТО. Третий сценарий включал совместные экономические и антитеррористические программы, признание границ Украины и России. Четвертый сценарий допускал интеграцию Украины в НАТО [Reynolds, Jensen 2024].

Полученные результаты позволяют сказать, что корректная постановка задач и встраивание обучающих примеров существенно повышают качество вывода модели и его соответствие текущим условиям конфликта. Более того, ИИ может быть использован для тестирования политических альтернатив и разработки сценариев, однако требует постоянной адаптации исходных данных, экспертной валидации и стратегической интерпретации полученных результатов.

ИИ может использоваться в качестве аналитического инструмента при моделировании дипломатических переговоров. Он особенно полезен на стадии формирования альтернатив, струк-

турирования аргументов и выявления повторяющихся элементов в исторических данных. В то же время, ключевым ограничением остается неспособность ИИ в полной мере учитывать быстро меняющуюся политическую динамику, стратегические интересы сторон и эмоциональные компоненты, присущие человеческому принятию решений. Следовательно, ИИ следует рассматривать не как автономный механизм выработки политики, а как вспомогательный аналитический ресурс, требующий критической настройки, контроля и интеграции с экспертным знанием.

Риски автоматизированных решений в ядерной области. Одним из наиболее обсуждаемых направлений современных исследований на пересечении ИИ и международной безопасности является вопрос о влиянии алгоритмических систем на стратегическую стабильность в условиях потенциальных вооруженных конфликтов, особенно между ядерными державами. Это направление объединяет вопросы кризисного принятия решений, теории сдерживания и когнитивной психологии. В данной области наблюдается быстрый рост интереса среди аналитиков, политологов и военных стратегов, однако академических публикаций с открытыми данными и четко обозначенной методологией почти нет.

Центральный вопрос в исследованиях – может ли автоматизация аналитики и решений привести к непреднамеренной эскалации. Вводится понятие «алгоритмической стабильности» – способности сохранять предсказуемость стратегических взаимодействий в условиях высокой автоматизации.

Ключевым вкладом в развитие данной темы является проект Algorithmic Stability, реализованный экспертами США. Его цель состояла в проверке гипотез о том, как ИИ могут изменить принципы управления кризисами, связанными с риском ядерной эскалации. Методом служили серия симуляций гипотетических международных кризисов, в которых моделировалось поведение государств с различным уровнем интеграции ИИ в структуру принятия решений. Исследование было построено на ряде допущений из теории международных отношений. Во-первых, эксперты опирались на теорию рационального выбора в условиях неопределенности, в рамках которой алгоритмы ИИ выполняют функцию усиления (но не замены) экспертизы человеком. Во-вторых, использовали концепт алгоритмической стабильности – термин, обозначающий способность системы стратегического взаимодействия сохранять предсказуемость и контролируемость при высокой степени автоматизации аналитики и решений.

Проект показал, что государства пока еще не передают критически важные решения полностью ИИ, но используют его для

усиления аналитики. Исследование показало необходимость прозрачности и новых режимов контроля над вооружениями, включая взаимный контроль за применением ИИ в стратегических системах. Более того, исследование позволило сделать два ключевых вывода. Во-первых, несмотря на различия в уровне технологической оснащенности разных государств, фундаментальные принципы стратегического поведения в условиях кризиса сохраняются. Государства продолжают применять комплексный подход, комбинируя инструменты давления — от дипломатии до экономических и информационных мер — и не склонны передавать критически важные решения полностью на усмотрение алгоритмических систем. Во-вторых, по мере развития ИИ-систем меняются условия выбора тактических и оперативных решений в конфликтных ситуациях. Это требует развития новых механизмов контроля над вооружениями, включая взаимную транспарентность в вопросе использования ИИ в критических системах⁴.

Дополнительную перспективу дает исследование о восприятии ИИ в стратегическом мышлении России и США. США, как отмечается, рассматривают ИИ преимущественно как инструмент повышения эффективности — автоматизации обработки данных, усиления разведки, логистики и поддержки принятия решений. Такие инициативы, как Project Maven и Центр ИИ при Министерстве обороны (JAIC), служат иллюстрацией этого подхода. В то же время Россия склонна использовать ИИ как инструмент специального назначения, делая акцент на автономных платформах, радиоэлектронной борьбе, а также кибер- и психологических операциях, включающих элементы ИИ-моделирования поведения противника. В итоге, специалисты приходят к выводу, что ИИ пока не угрожает стратегической стабильности. Его влияние зависит от институционального контекста, степени прозрачности, политической воли и способности государств к взаимному доверию и диалогу. При отсутствии этих условий возникает угроза недопонимания, автоматической эскалации и нарушения баланса сдерживания [Nadibaidze, Miotto 2023].

Процесс принятия внешнеполитических решений. ИИ оказался эффективным инструментом для анализа внешнеполитических решений лидеров на основе публичных речей, интервью, стратегических документов. Все модели ИИ позволяют выявлять

⁴ Jensen B., Atalan Ya., Macias III J.M. Algorithmic stability: How AI could shape the future of deterrence // Center for Strategic and International Studies. 2024. URL: <https://www.csis.org/analysis/algorithmic-stability-how-ai-could-shape-future-deterrence> (дата обращения 04.03.2025).

эмоциональные паттерны, риторические стратегии, исторические аналогии и влияние советников на процесс принятия решений, что существенно трансформирует такую область знаний, как внешнеполитический анализ. Сегодня можно выделить несколько перспективных направлений исследований.

Во-первых, особое внимание уделяется текстовому анализу, с чем ИИ справляется особенно успешно. Модели способны отслеживать как прямые упоминания, так и скрытые параллели – структурные совпадения между прошлыми и текущими ситуациями. По мере роста использования исторических аналогий современными лидерами возрастает интерес к выявлению их влияния на решения. Специалисты в области внешнеполитического анализа уже давно выдвинули гипотезу, что использование исторических аналогий определяет содержание будущего решения. Исследования показывают, что аналогии выступают как средство обоснования решения, а также как способ справиться с психологическим давлением и стрессом, связанным с весом принимаемых решений [Khong 1992; Jervis 2017; Цветкова 2024]. Однако на протяжении многих лет эксперты не имели возможности провести масштабный анализ источников и сопоставить, на каких этапах и в какой мере лидеры различных стран используют исторические аналогии и насколько серьезно они влияют на итоговое решение.

Открытые ИИ-инструменты, такие как ChatGPT, как оказалось, могут быть эффективно использованы для поиска и систематизации исторических аналогий и анализа процессов принятия внешнеполитических решений в общем. Ключевым элементом выступает возможность загрузки больших объемов информации – заявлений и интервью президентов и их советников, а также доктринальных документов – и проведения целенаправленного анализа на предмет наличия поиска позиций, эмоций, предпочтений президента, исторических аналогий, влияния советников и пр.

Благодаря точности в поиске аналогий – как явных (например, прямые сравнения отдельных кризисов с событиями из прошлого, такими как Мюнхен, Ялта или Карибский кризис), так и имплицитных (сопоставление по структуре аргументации или набору исходных характеристик) – становится возможным не только фиксировать факт их появления, но и определять их функцию: когнитивную (влияние на ход размышления) или риторическую (убеждение аудитории).

Последние исследования показывают, что исторические аналогии рассматриваются не как вспомогательный аргумент, сопровождающий процесс выбора решения среди предлагаемых альтернатив, а как индикатор уже принятого решения лидером. Возникающая

в публичной риторике историческая аналогия служит предвестником действия, а не реакцией на него. Эмпирический материал, включающий кейсы лидеров самых разных стран, подтверждает, что артикуляция аналогии в риторике лидеров предшествует ключевым внешнеполитическим шагам. Последовательность уточняющих запросов к модели позволяет реконструировать логику принятия решений, идентифицируя не только момент появления аналогии, но и ее аргументационную структуру. В отличие от классических работ, где аналогии трактуются как инструмент убеждения, здесь они интерпретируются как маркеры хода мысли национальных лидеров и их советников, особенно в условиях стратегической неопределенности. ИИ, при правильной настройке и контекстной подаче информации, способен воспроизвести и классифицировать аналогии с учетом международного контекста, выделяя те, которые имеют значение для прогнозирования [Tsvetkova 2025].

Вторым интересным трендом стали исследования влияния советников на позицию лидеров государства, в частности президентов США. Методы машинного анализа позволяют сопоставлять мнения советников и президентов, а также отслеживать идеологическую направленность решений, что приводит исследователей к выводу, что лидеры не всегда следуют мнению большинства – они могут целенаправленно усиливать влияние отдельных фигур, исходя из собственных установок. К данному направлению относятся исследования, посвященные влиянию советников на президента в процессе внешнеполитического принятия решений. Модели искусственного интеллекта позволяют собрать и систематизировать новый массив данных о персональном составе и риторике советников, по-новому реконструировать протоколы заседаний, в том числе Совета национальной безопасности США (СНБ), а также выявить степень зависимости позиции президента от конкретных фигур в его окружении. Кроме того, ИИ-инструменты позволяют вычислить долю решений, принятых в агрессивной или, напротив, миролюбивой логике.

Анализ показывает, что ИИ-модели формируют более сложное представление о механизме предпочтений главы государства. Согласно предложенной модели, президент может намеренно усиливать влияние отдельных советников в зависимости от собственных политических установок, идеологических ориентиров или эмоционального состояния. Авторы исследования опираются на широкий спектр источников: рассекреченные протоколы СНБ, мемуары, интервью, а также документы, содержащие персонализированные портреты советников. Особое внимание уделяется машинному анализу совпадения позиций президента и его окружения в конкрет-

ных кейсах, а также оценке идеологических установок советников как ключевого фактора влияния на решения главы государства [Hafner-Burton, Haggard, Victor 2024].

В-третьих, ИИ-инструменты применяются для выявления векторов внешнеполитической ориентации (на основе речей в ООН, официальных выступлений), анализа стратегических нарративов (особенно в Китае, где США воспринимаются одновременно как угроза и ориентир), а также для мониторинга изменений в дипломатическом дискурсе. Применение методов ИИ к корпусам речей и голосований в Генеральной Ассамблее ООН позволяет выявлять, например, «вектор политической близости» между странами. Исследования показывают, что несмотря на развитие экономических связей с Китаем, ряд стран Центральной и Восточной Европы сохраняют риторическую и политическую ориентацию на США и Германию. Эти данные ставят под сомнение распространенные представления о геополитическом «повороте к Китаю» и подтверждают, что тексты дипломатических выступлений могут выступать валидным индикатором внешнеполитической ориентации государств [Turcsanyi, Liškutin, Mochtak 2023].

Наконец, значимый вектор исследований – анализ национальных стратегических нарративов с помощью автоматизированных методов. В частности, машинный анализ правительственных и экспертных оценок по тематике военного потенциала показывает, что США играют роль центрального ориентира в китайской экспертизе. Автоматизированный поиск ключевых слов, дискурсивных шаблонов и тематических фреймов позволяет установить, что китайская стратегическая мысль выстраивает свои цели через постоянное сопоставление с американскими достижениями в военной сфере. Восприятие США как одновременно угрозы и примера для подражания оказывает воздействие на стратегическую повестку КНР. Значимым вкладом является формирование специализированных баз данных по пресс-конференциям или высказываниям лидеров в социальных сетях, позволяющих проводить анализ внешнеполитической риторики в динамике и отслеживать эволюцию политических нарративов, изменений в стиле и содержании публичных выступлений [Mochtak, Turcsanyi 2021].

Заключение

Применение искусственного интеллекта в анализе внешней политики и международных отношений постепенно становится неотъемлемой частью как академических исследований, так и прак-

тической аналитики. Модели ИИ уже сегодня используются для прогнозирования международных кризисов, тестирования дипломатических сценариев, оценки рисков эскалации и реконструкции логики принятия решений лидерами государств. Наибольший эффект достигается при использовании локальных и специализированных моделей, настроенных на конкретные задачи и обученных на релевантных корпусах документов.

Опыт США, Китая и России показывает, что архитектура моделей, принципы обучения и сферы применения ИИ существенно различаются, отражая особенности стратегической культуры, уровень технологической зрелости и институциональные приоритеты. Однако вне зависимости от контекста, ключевым звеном в работе с ИИ остается эксперт: именно он интерпретирует выводы модели, выявляет риски предвзятости, уточняет параметры и обеспечивает критическую верификацию результатов.

ИИ в международных исследованиях – это не самостоятельный субъект принятия решений, а инструмент аналитического сопровождения, расширяющий когнитивные возможности исследователя, и позволяющий выявлять скрытые паттерны, предсказывать альтернативные сценарии и моделировать поведение государств и политических лидеров в условиях стратегической неопределенности. Будущие исследования должны быть направлены на адаптацию моделей к локальному контексту и интеграцию ИИ в гибридные экспертные системы принятия решений. Специализированные ИИ, которые ограничены локальным корпусом источников, создают предпосылки для более точного прогнозирования и анализа ситуации.

Литература

- Цветкова 2020 – *Цветкова Н.А.* Феномен цифровой дипломатии в международных отношениях и методология его изучения // Вестник РГГУ. Серия: «Политология. История. Международные отношения». 2020. № 2. С. 37–47.
- Цветкова 2024 – *Цветкова Н.А.* Роль исторических аналогий в процессе принятия внешнеполитических решений в США, 1945–2024 // Американский ежегодник – 2024. М.: Весь мир, 2024. С. 57–80.
- Hafner-Burton, Haggard, Victor 2024 – *Hafner-Burton E.M., Haggard S., Victor D.G.* Advisers and aggregation in foreign policy decision-making // American Political Science Review. 2024. Vol. 118. No. 2. P. 435–450. DOI: 10.1017/S0003055424000081.
- Jensen, Reynolds, Atalan 2024 – *Jensen B., Reynolds I., Atalan Ya.* What AI thinks about foreign policy decisions: Benchmarking models for strategy and statecraft //

- CSIS. 2024. URL: <https://www.csis.org/analysis/what-ai-thinks-about-foreign-policy-decisions-benchmarking-models-strategy-and-statecraft> (дата обращения 04.03.2025).
- Jervis 2017 – *Jervis R.* Perception and misperception in international politics. Princeton: Princeton University Press, 2017. 544 p.
- Khong 1992 – *Khong Yu.F.* Analogies at war: Korea, Munich, Dien Bien Phu, and the Vietnam decisions of 1965. Princeton: Princeton University Press, 1992. 286 p.
- Mochtak, Turcsanyi 2021 – *Mochtak M., Turcsanyi R.Q.* Studying Chinese foreign policy narratives: Introducing the Ministry of Foreign Affairs press conferences corpus // *Journal of Chinese Political Science*. 2021. Vol. 26. P. 743–761. DOI: 10.1007/s11366-021-09762-3.
- Nadibaidze, Miotto 2023 – *Nadibaidze A., Miotto N.* The impact of AI on strategic stability is what states make of it: comparing US and Russian discourses // *Journal for Peace and Nuclear Disarmament*. Vol. 6. No. 1. P. 47–67.
- Reynolds, Jensen 2024 – *Reynolds I., Jensen B.* Can AI help end a war? Using generative AI to explore peace scenarios in Ukraine // CSIS. 2024. URL: <https://www.csis.org/analysis/can-ai-help-end-war-using-generative-ai-explore-peace-scenarios-ukraine> (дата обращения 04.03.2025).
- Tsvetkova 2025 – *Tsvetkova N.A.* The predictive power of historical analogies in foreign policy decision-making: An AI-assisted approach [Preprint] // *Humanities and Social Sciences Communications*. Springer Nature/Research Square, 2025. URL: <https://www.researchsquare.com/article/rs-6490065> (дата обращения 04.03.2025).
- Turcsanyi, Liškutin, Mochtak 2023 – *Turcsanyi R.Q., Liškutin K., Mochtak M.* Diffusion of influence? Detecting China's footprint in foreign policies of other countries // *Chinese Political Science Review*. 2023. Vol. 8. P. 461–486. DOI: 10.1007/s41111-022-00217-5.

References

- Nadibaidze, A. and Miotto, N. (2023), “The impact of AI on strategic stability is what states make of it: comparing US and Russian discourses”, *Journal for Peace and Nuclear Disarmament*, vol. 6, no. 1, pp. 47–67.
- Hafner-Burton, E.M., Haggard, S. and Victor, D.G. (2024), “Advisers and aggregation in foreign policy decision-making”, *American Political Science Review*, vol. 118, no. 2, pp. 435–450, <https://doi.org/10.1017/S0003055424000081>.
- Jensen, B., Reynolds, I. and Atalan, Ya. (2024), “What AI thinks about foreign policy decisions: Benchmarking models for strategy and statecraft”, *Center for Strategic and International Studies*, available at: <https://www.csis.org/analysis/what-ai-thinks-about-foreign-policy-decisions-benchmarking-models-strategy-and-statecraft> (Accessed 4 March 2025).
- Jervis, R. (2017), *Perception and misperception in international politics*, Princeton University Press, Princeton, USA.

- Khong, Yu.F. (1992), *Analogies at war: Korea, Munich, Dien Bien Phu, and the Vietnam decisions of 1965*, Princeton University Press, Princeton, USA.
- Mochtak, M. and Turcsanyi, R.Q. (2021), "Studying Chinese foreign policy narratives: Introducing the Ministry of Foreign Affairs press conferences corpus", *Journal of Chinese Political Science*, vol. 26, pp. 743–761, <https://doi.org/10.1007/s11366-021-09762-3>.
- Reynolds, I. and Jensen, B. (2024), "Can AI help end a war? Using generative AI to explore peace scenarios in Ukraine", *Center for Strategic and International Studies*, available at: <https://www.csis.org/analysis/can-ai-help-end-war-using-generative-ai-explore-peace-scenarios-ukraine> (Accessed 4 March 2025).
- Tsvetkova, N.A. (2020), "The digital diplomacy as a phenomenon of international relations: Research methodology", *RSUH/RGGU Bulletin. "Political Science. History. International Relations" Series*, no. 2, pp. 37–47.
- Tsvetkova, N.A. (2024), "The role of historical analogies in U.S. foreign policy decision-making, 1945–2024", *Amerikanskii ezhegodnik – 2024* [American Yearbook – 2024], Ves' mir, Moscow, Russia, pp. 57–80.
- Tsvetkova, N.A. (2025), "The predictive power of historical analogies in foreign policy decision-making: An AI-assisted approach" [Preprint], *Humanities and Social Sciences Communications*, Springer Nature/Research Square, available at: <https://www.researchsquare.com/article/rs-6490065> (Accessed 4 March 2025).
- Turcsanyi, R.Q., Liškutin, K. and Mochtak, M. (2023), "Diffusion of influence? Detecting China's footprint in foreign policies of other countries", *Chinese Political Science Review*, vol. 8, pp. 461–486, <https://doi.org/10.1007/s41111-022-00217-5>.

Информация об авторах

Наталья А. Цветкова, доктор исторических наук, профессор, Институт США и Канады РАН им. Г.А. Арбатова, Москва, Россия; 117997, Россия, Москва, Хлебный пер., д. 2/3, стр. 1; n.tsvetkova@iskran.ru
ORCID ID: 0000-0002-0156-9069

Information about the author

Natalia A. Tsvetkova, Dr. of Sci. (History), professor, Georgy Arbatov Institute for US and Canada Studies, Russian Academy of Sciences, Moscow, Russia; bldg. 1, bld. 2/3, Khleby Line, Moscow, Russia, 117997; n.tsvetkova@iskran.ru

ORCID ID: 0000-0002-0156-9069