

Бутко В. Н.,

кандидат технических наук, доцент,
marnic08@yandex.kz

*Костанайский социально-технический университет
имени академика З. Алдамжар,
110000 г. Костанай, пр-т Кобыланды Батыра, 27*

ИСТОРИЯ СТАНОВЛЕНИЯ И ПРОБЛЕМНАЯ СУЩНОСТЬ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

Аннотация: *В настоящее время Человечество находится на перепутье. Родившаяся в эпоху генеративного искусственного интеллекта (ИИ) и набирающая всё большую силу новая технология пространственного интеллекта - взаимодействие систем ИИ с трехмерным виртуальным и физическим миром - может принести человечеству не только пользу, но и новые угрозы. Особого внимания заслуживают такие негативные стороны ИИ, как алгоритмическая предвзятость, дезинформация и использование автономных систем. Такие проблемы можно преодолеть только через глобальное сотрудничество. А для возможности рационального понимания и оптимального влияния на развитие этого процесса представляется необходимым вначале исследовать историю становления и сущность искусственного интеллекта (ИИ).*

Планируемая основа данного исследования - попытаться выяснить историю становления и сущность искусственного интеллекта. В работе используются следующие научные методы исследования: аналитический, исторический, логический, экспертных оценок.

Применение искусственного интеллекта расширяется и ускоряется быстрыми темпами в глобальном масштабе. Однако этот технологический прогресс может обернуться как победой, так и катастрофой, формируя глобальный вызов человечеству. Действительно, ИИ высшего уровня способен самостоятельно изменять свои код и цель оптимизации. А его возможности всё более быстрого интеллектуального самосовершенствования вызывает значительное возрастание глобальных рисков. Но даже если ИИ будет развиваться только за счёт внешних факторов (исключительно силами людей-инженеров и учёных), его интеллектуальная эффективность также может сделать резкий скачок. При этом - никто не будет никого предупреждать о неожиданном моменте «рождения» самодостаточного ИИ, уже полностью независимого от Человека.

И самое трагичное – всё это происходит в условиях, когда интеллект подавляющей части Человечества, способный устранить эти риски, очень

старательно пытаются ограничить, «обнулить» современные Глобальные Сатанисты, которые стремятся к абсолютному управлению миром - «элита» англо-саксов и, прежде всего - управляющая «элита» Новой Британской империи. Они разрушают культуру, образование и науку всех уровней путём развязывания терроризма, войн, пропаганды неонацизма, ЛГБТ, внедрения пресловутой Болонской системы в образовании и т.д. и т.п.

Мы до сих пор не до конца понимаем, как ведут себя продвинутые системы ИИ. Именно поэтому большинство экспертов согласно, что мировому сообществу необходимо эффективное управление ИИ. Однако здесь возникает новая проблема – каким мировым силам удастся первыми (а вторых уже не будет!) поставить под контроль себе искусственный интеллект – Силам Зла, которые сейчас превалируют в мире, или Силам Добра, которые сейчас только пытаются перехватить «Руль управления земным Миром» у Сил Зла. Поэтому необходимы дальнейшие тщательные исследования ИИ, как фактора, прежде всего, Глобального Управления.

***Ключевые слова:** история ИИ, сущность, негативные стороны, глобальное сотрудничество, код и цель ИИ, самосовершенствование ИИ, глобальные риски, самодостаточность, управление ИИ, контроль.*

*«Я пришел к убеждению,
что искусственный интеллект
и связанные с ним дисциплины
изменяют человеческое сознание,
как это произошло в эпоху Просвещения»
Генри Киссинджер [1].*

Введение

Профессор Стэнфордского института искусственного интеллекта Фэй-Фэй Ли более 25 лет изучает и разрабатывает технологии искусственного интеллекта (ИИ). Она считает, что *в настоящее время Человечество находится на перепутье*. Никогда ещё в своей истории оно не находилось «на таком уникальном пересечении научных возможностей и глобальной ответственности». В эпоху генеративного ИИ, исследуя новую технологию - пространственный интеллект – взаимодействие систем ИИ с трехмерным виртуальным и физическим миром, эксперт делает вывод, что *эта технология может принести человечеству не только пользу, но и новые угрозы*. По её мнению, особого внимания заслуживают такие негативные стороны ИИ, как алгоритмическая предвзятость, дезинформация и использование автономных систем. «Но нам многое еще предстоит понять о масштабах этих опасных явлений и их последствиях», – отмечает она. И добавляет, что эти проблемы можно преодолеть только через глобальное сотрудничество [2]. А для возможности рационального понимания и оптимального влияния на развитие этого процесса представляется необходи-

мым вначале исследовать историю становления и сущность искусственного интеллекта (ИИ).

Методология

Для выяснения истории становления и сущности искусственного интеллекта в настоящей работе используются следующие научные методы исследования: аналитический, исторический, логический, экспертных оценок, стеганоанализ.

Обзор литературы, результаты

1. История становления искусственного интеллекта (ИИ). Первые объекты ИИ (компьютеры) появились в середине XX века [3]. Дальнейшее его развитие характеризуется несколькими последовательными этапами.

Первый (1950-е(40-е) — 1970-е годы) - зарождение искусственного интеллекта как науки и технологии. *По одной из версий*, отправным моментом является летний исследовательский проект Дартмутского колледжа по ИИ 1956 года. Тогда, в рамках состоявшегося семинара, был осуществлён сеанс мозгового штурма, ставший ключевым событием в области ИИ. После ряда взлётов и падений, постепенно происходил рост влияния ИИ на общественное сознание. *По другой версии (оборонной)* – начало развития ИИ – 1943 год. В этом году был создан *первый электронный компьютер ENIAC (ЭНИАК)*. Он использовался для расчетов траекторий артиллерийских снарядов и других военных задач [3,4].

- В 1950-1960-х годах были разработаны различные системы автоматического управления оружием, основанные на ИИ, такие как система SAGE (Semi-Automatic Ground Environment) для обнаружения и перехвата воздушных целей, система Distant Early Warning Line (DEW, Дью) для предупреждения о ракетных атаках.

- В 1960-е годы США и СССР разрабатывали *системы раннего предупреждения о ракетном нападении* с использованием ИИ для обработки данных с радаров и спутников.

- В 1970-е годы Израиль использовал ИИ для *анализа фотографий с разведывательных самолетов* и определения местонахождения египетских войск перед войной Йом-Киппур [3].

Второй этап развития ИИ (1980-е — 1990-е годы) - расцвет и кризис.

- В 1980-е годы США использовали ИИ для создания системы "Звездных войн", которая предполагала *использование космических оружейных платформ* для перехвата советских баллистических ракет [3].

- В 1985 г. один из создателей ИИ Рэй Соломонофф ввел понятие «точки бесконечности» во временной шкале ИИ, а также проанализировал величину «будущего шока», который «мы можем ожидать от нашего расширенного научного сообщества ИИ», а также социальные эффекты.

- В 1988 г. ученые Холлоуэй и Хэнд опубликовали статью об ИИ, которая начиналась следующим образом: «Искусственный интеллект — это больше не академический термин, а реальность» [4].

Третий этап (2000-е — настоящее время) - возрождение и революция ИИ. Произошел резкий скачок в развитии ИИ - появились мощные компьютеры, большие объемы данных, новые алгоритмы и методы машинного обучения. Были созданы ИИ-системы, способные обыгрывать человека в сложные игры, распознавать лица и объекты на изображениях, генерировать тексты и изображения, управлять автомобилями и дронами. Постепенно расширялся интерес к ИИ со стороны науки, образования, бизнеса, правительства и общественности [3].

Сейчас использование технологий ИИ все больше совершенствуется и набирает обороты. По прогнозу исследовательской компании *Gartner*, в 2022 г. объемы мирового рынка программного обеспечения, использующего алгоритмы ИИ, должны были составить 62,5 млрд. долл., что на 21,3% больше, чем в 2021 г. Прогноз оправдался и объем мирового рынка программного обеспечения с технологиями ИИ по итогу 2022 г. составил 64 млрд долл [1].

2. Проблемная сущность искусственного интеллекта. ИИ — это общий термин, относящийся к любому типу компьютерного программного обеспечения, участвующего в обучении человека, планировании и поиске различных вариантов решений. Это новая сфера развития человечества. ИИ позволяет машинам учиться на опыте, приспосабливаться к новым данным и выполнять задачи вместо человека. Быстрота и высокая эффективность процессов вызывают серьезные опасения относительно утраты способности человека к самостоятельному мышлению, формированию собственного восприятия мира, взаимоотношению с другими людьми и т.п. [1,5].

Генеральный директор Alphabet Group *Сундар Пичаи* считает, что ИИ является величайшим открытием человечества, на уровне огня и электричества, и изменит функционирование общества. По его мнению, создание ИИ "так же фундаментально, как создание микропроцессора, персонального компьютера, Интернета и мобильного телефона".

Сэм Альтман, генеральный директор OpenAI, считает, что ИИ благодаря своей универсальной природе обладает безграничным потенциалом и способен выполнять практически любые народнохозяйственные задачи. Однако нерегулируемые разработки ИИ несут риски для конфиденциальности и безопасности [6].

ИИ, в отличие от естественного, которым владеют люди и животные, выражается в интеллектуальном поведении или деятельности машины. В зависимости от его способностей мыслить, делать умозаключения и решать проблемы, *он делится на слабый и сильный. Слабый ИИ* — это системы, которые способны выполнять определенные задачи лучше или быстрее, чем человек, но не обладают самосознанием или общим интеллектом. *Сильный* - моделирует творческую личность, обладает мышлением, восприятием и даже самосознанием, но достигнутый им уровень к настоящему времени пока не ясен. Однако есть опасения, что в будущем ИИ превысит уровень человеческого со всеми вытекающими отсюда последствиями, вплоть до управления самим человечеством [7,3,1].

*Проблемность восприятия и плодотворного анализа ИИ состоит не в его сложности, а, прежде всего, в антропоморфическом его представлении. Дело в том, что в мирозданческой действительности человеческий интеллект – это лишь микроскопическая точка в океане возможных устройств ума. Это одна из главных причин, по которой людям кажется, что они знают об Искусственном Интеллекте гораздо больше, чем это есть на самом деле. ИИ – это **передовой, инновационный процесс, а не результат, не твёрдо установленное знание**. В связи с этим, представляется довольно странным, что проблема глобальных рисков в связи с искусственным интеллектом фактически не обсуждалась открыто, широко в существующей технической литературе до самого недавнего времени (по крайней мере, по состоянию до января 2006 года). Кроме того, антропоморфизм в процессе исследования и оптимизации использования ИИ является не рациональным и мало перспективным способом мышления [8].*

Особое место занимает развитие так называемого *генеративного искусственного интеллекта*, который способен генерировать различные информационные данные. Правда, тут стоит также упомянуть, что способен он это делать при наличии соответствующих подсказок со стороны живого человека [9]. ***Наша неспособность понять, как "думает" ИИ, и ограниченное понимание вторичных и третичных эффектов наших команд или запросов к ИИ также вызывают серьезное беспокойство.*** А если учесть, что даже людям нередко достаточно сложно взаимодействовать между собой, то как мы будем управлять нашими отношениями ещё дополнительно с одним или несколькими ИИ, интеллектуально и технологически превосходящими человека? [10]. Тем более, когда ИИ сможет самостоятельно переписать свою существующую когнитивную функцию, значительно улучшив её работу [8]. Таким образом, ИИ может выполнять действия, подобные человеческим, такие как вождение, производство и другие физические задачи. Несомненно, этот потенциал ИИ привлекает повышенное внимание различных объектов и субъектов. Например, согласно The Wall Street Journal «30 процентов организаций тестируют проекты ИИ. Около 47 процентов использовали ИИ в своих бизнес-процессах» [4].

3. Современный уровень ИИ и его перспективы. Стремительное развитие ИИ до третьего поколения - искусственного суперинтеллекта (ИСИ) резко повышает риски выхода ИСИ из-под контроля человека. Действительно, по мнению учёных, «технически даже стандартный микропроцессор работает в 10 миллионов раз быстрее, чем человеческий нейрон, и компьютеры могут запомнить больше информации за 1 секунду, чем человек мог бы запомнить за всю жизнь» [4]. Поэтому *ИСИ легко сможет превзойти людей*. К тому же, если люди привыкли думать на человеческом уровне, то система ИСИ будет думать на уровне ИСИ. А по мнению ученых, «так же, как люди никогда не смогут по-настоящему понять, как думают шимпанзе несмотря на то, что они разделяют 99% нашей ДНК, мы не сможем понять, как думает система ИСИ. Это ограничивает нашу способность контролировать такие системы, что снова делает их более опасными, чем полезными» [4].

Если мы считаем конечной целью замену, в той или иной степени, человеческого интеллекта машинным, то сразу возникает масса этических и правовых вопросов в отношении ИСИ. Здесь *главная дилемма* - передача технологиям человеческих когнитивных функций и рисков с этим связанных. Теперь не только человек, но и машина начинают приобретать привелегии в обучении, решении проблем и мышлении [4]. *Мало того, быстрая эволюция ИИ делает его непредсказуемым. Даже технологические лидеры, принимающие активное участие в развитии этой технологии, не могут предугадать новые области применения, проблемы и прорывы. Автономная природа ИИ позволяет ему развиваться самостоятельно. При получении достаточного количества данных, системы ИИ продолжают учиться и совершенствоваться* [6].

Реакция бизнес-лидеров на продукты с искусственным интеллектом, такие как ChatGPT, была неоднозначной - от волнения до видений конца света. Например, такие, как Билл Гейтс (американский бизнес-магнат и соучредитель Microsoft), считают, что ИИ может "изменить наш мир", сделав работу более эффективной. Инструменты типа ChatGPT могут высвободить время в жизни работников, повышая эффективность сотрудников, станут "величайшей вещью в этом десятилетии" [6]. Другие, например Илон Маск, Генеральный директор Tesla и SpaceX, технический директор X и владелец недавно созданного стартапа в области искусственного интеллекта XAI, заявил, что *ИИ будет иметь потенциал стать «самой разрушительной силой в истории». «У нас будет нечто, что впервые будет умнее самого умного человека»* [11]. Илон Маск считает, что ИИ – это «один из самых больших рисков для будущего цивилизации» [6]. В целом, *Илон Маск* считает перспективы развития ИИ неоднозначными: «Трудно сказать точно, что это за момент, но наступит момент, когда люди уже не понадобятся для работы. Вы сможете найти работу, если хотите иметь работу для личного удовлетворения. Но ИИ сможет сделать всё. Я не знаю, насколько комфортно это будет людям. Это как найти волшебного джинна, который исполнит любое ваше желание, при этом желаний может быть бесконечно много». *И. Маск* неоднократно предупреждал об угрозах, которые ИИ представляет для человечества, считая его *более опасным, чем ядерное оружие* [11].

Заключение

Применение искусственного интеллекта расширяется и ускоряется быстрыми темпами в глобальном масштабе, становясь все более многосторонним, затрагивая области науки и общественной жизни, значительно усиливая свою роль в мире. Однако этот технологический прогресс может обернуться как победой, так и катастрофой, формируя глобальный вызов человечеству. Самое страшное состоит в том, что, учитывая скорость и масштабы развития ИИ – он делает только свои первые шаги [12,4]. К тому же, ИИ высшего уровня способен изменить свои код и цель оптимизации. А его возможности всё более быстрого интеллектуального самосовершенствования, вызывает значительное возрастание глобальных рисков. Но даже если ИИ будет развиваться только за счёт внешних факторов (исключительно силами людей-инженеров), его интеллекту-

альная эффективность также может сделать резкий скачок. При этом - никто не будет никого предупреждать о неожиданном моменте «рождения» самодостаточного ИИ, уже полностью независимого от человека [8].

И самое трагичное – всё это происходит в условиях, когда интеллект подавляющей части Человечества, способный устранить эти риски, очень старательно пытаются ограничить, «обнулить» современные Глобальные Сатанисты, которые стремятся к абсолютному управлению миром - «элита» англо-саксов и, прежде всего - управляющая «элита» Новой Британской империи, разрушая культуру, образование и науку всех уровней путём развязывания терроризма, войн, пропаганды неонацизма, ЛГБТ, внедрения пресловутой Болонской системы в образовании и т.д. и т.п.

Поскольку развитие ИИ является необратимым, это должно стать новой глобальной проблемой для политиков, лидеров бизнеса, исследователей ИИ и общественности: какие рамки необходимо создать для совершенствования моральных принципов развития ИИ и предотвращения его потенциальных негативных последствий. На сегодняшний день ИИ все еще находится на начальной стадии развития, поэтому необходимо внимательно наблюдать и осматривательно подходить к вопросам его практического применения [7, с. 200]. *Мы до сих пор не до конца понимаем, как ведут себя продвинутые системы ИИ. Такое отсутствие прозрачности, доступа, вычислительных и других ресурсов, а также понимания препятствует возможности определить то, откуда возникают риски и на ком должна лежать ответственность за управление этими рисками (или компенсацию ущерба)* [13]. При этом, злонамеренное использование и безответственное развитие ИИ несет в себе угрозы национальной и международной безопасности. Именно поэтому большинство экспертов согласно с тем, что мировому сообществу необходимо эффективное управление ИИ [14]. Дело в том, что «ИИ играет все более важную роль в различных сферах нашей жизни. Однако его влияние не ограничивается только отдельными областями, и он может иметь значительные последствия для глобального порядка и международных отношений [3].

Таким образом, ИИ ворвался во многие сферы человеческой жизни и явно вышел на глобальный уровень обсуждения, принятия или не принятия последствий ИИ. Коснулся он и международных организаций в контексте развития глобальных процессов [4]. Однако здесь возникает новая проблема – каким мировым силам удастся первыми (а вторых уже не будет!) поставить под контроль себе искусственный интеллект – Силам Зла, которые сейчас превалируют в мире, или Силам Добра, которые сейчас только пытаются перехватить «Руль управления земным Миром» у Сил Зла. Поэтому необходимы дальнейшие тщательные исследования ИИ, как фактора, прежде всего, Глобального Управления.

ЛИТЕРАТУРА

1. Андрей Волков - Искусственный интеллект и международные отношения: социальные аспекты и влияние на международную безопасность // 31 января 2024 - <https://russiancouncil.ru/blogs/a-volkov/iskusstvennyy-intellekt-i-mezhdunarodnye-otnosheniya-sotsialnye-aspekt/> (Дата обращения 04.03.2025).
2. Глава ООН – о будущем ИИ: прогресс и новые вызовы для человечества //Новости ООН. 19 декабря 2024. - <https://news.un.org/ru/> (Дата обращения 12.06.2025).
3. Новый мировой порядок: роль искусственного интеллекта в формировании глобальной политики. //09:35 / 10 августа, 2023 - <https://www.securitylab.ru/analytics/540782.php> (Дата обращения 02.03.2025).
4. Загайнов Михаил Рудольфович - Искусственный интеллект – на пути к урегулированию в глобальном мире //Текст научной статьи по специальности «Право» - <https://cyberleninka.ru/article/n/iskusstvennyy-intellekt-na-puti-k-uregulirovaniyu-v-globalnom-mire> (Дата обращения 04.03.2025).
5. Инициатива глобального управления искусственным интеллектом // МИД КНР. - 2023-10-20 16:05 - https://www.fmprc.gov.cn/rus/wjdt/gb/202310/t20231024_11166700.html (Дата обращения 04.03.2025).
6. Герман Геншин - Вот что думают 10 лидеров мирового технологического рынка об искусственном интеллекте //15 декабря 2023 - <https://vc.ru/u/2830724-andrew-volpert/1004639-samy-e-vliyatelnye-lyudi-v-mire-ob-ii> (Дата обращения 04.03.2025).
7. Ван Цзяньган, кандидат исторических наук - Влияние развитие искусственного интеллекта на процессы глобализации // Теории и проблемы политических исследований. 2019. Том 8. № 1А. С. 193-202. - <http://www.publishing-vak.ru/file/archive-politology-2019-1/23-wang-jiangang.pdf> (Дата обращения 02.03.2025).
8. Э. Юлковский - Искусственный интеллект как позитивный и негативный фактор глобального риска //Биология, психология, медицина, демография и социология Э. Юлковский — Singularity Institute for Artificial Intelligence Palo Alto, CA Перевод: Алексей Валерьевич Турчин - <https://spkurdyumov.ru/biology/iskusstvennyj-intellekt/> (Дата обращения 04.03.2025).
9. Для чего глобалистам искусственный интеллект? //Еженедельник "Звезда", 22 января 2024. - <https://dzen.ru/a/Za41RPqCO1XCuvO3> (Дата обращения 02.03.2025).
10. Искусственный интеллект и геополитика. Как ИИ может повлиять на взлёт и падение наций? // 8 ноября 2023 - RAND Corporation Research <https://doi.org/10.7249/PEA3034-1> 2023 - <https://dzen.ru/a/ZUq-4lmlc2ikf2hq> (Дата обращения 04.03.2025).
11. Илон Маск: «Самая разрушительная сила в истории будет умнее самого умного человека и сможет выполнять всю работу», - <https://www.ixbt.com/news/2023/11/06/samaja-razrushitelnaja-sila-v-istorii-budet->

umnee-samogo-umnogo-cheloveka-i-smozhet-vypolnjat-vsju-rabotu--ilon-mask.html (Дата обращения 04.03.2025).

12. Илон Маск рассказывает о взрывном росте искусственного интеллекта и о том, что это значит для будущего // Brenda Kanana // Обновлено: 21 июня 2024 г. - <https://www.cryptopolitan.com/ru/elon-musk-reveals-ai-explosive-growth/> (Дата обращения 04.03.2025).

13. Управление ИИ в интересах человечества // AI Advisory Body Interim Report: Governing AI for humanity This translation is prepared by MGIMO AI Center. The present work is an unofficial translation for which the publisher accepts full responsibility. Неофициальный перевод документа выполнен сотрудниками Центра искусственного интеллекта МГИМО. (Промежуточный отчет). - Декабрь 2023. - <https://www.un.org/en/ai-advisory-body>. (Дата обращения 05.03.2025).

14. Васильев Д.П. Формирование международных режимов управления искусственным интеллектом: ключевые тенденции и основные акторы. - <https://cyberleninka.ru/article/n/formirovanie-mezhdunarodnyh-rezhimov-upravleniya-iskusstvennym-intellektom-klyuchevye-tendentsii-i-osnovnye-aktory> (Дата обращения 04.03.2025).

REFERENCES

1. Andrej Volkov - Iskusstvennyj intellekt i mezhdunarodnye otnosheniya: social'nye aspekty i vliyanie na mezhdunarodnuyu bezopasnost' // 31 yanvarya 2024 - <https://russiancouncil.ru/blogs/a-volkov/iskusstvennyy-intellekt-i-mezhdunarodnye-otnosheniya-sotsialnye-aspekt/> (Дата обращения 04.03.2025).

2. Glava OON – o budushchem II: progress i novye vyzovy dlya chelovechestva // Novosti OON. 19 dekabrya 2024. - <https://news.un.org/ru/> (Дата обращения 12.06.2025).

3. Novyj mirovoj poryadok: rol' iskusstvennogo intellekta v formirovanii global'noj politiki. // 09:35 / 10 avgusta, 2023 - <https://www.securitylab.ru/analytics/540782.php> (Дата обращения 02.03.2025).

4. Zagajnov Mihail Rudol'fovich - Iskusstvennyj intellekt – na puti k uregulirovaniyu v global'nom mire // Tekst nauchnoj stat'i po special'nosti «Pravo» - <https://cyberleninka.ru/article/n/iskusstvennyy-intellekt-na-puti-k-uregulirovaniyu-v-globalnom-mire> (Дата обращения 04.03.2025).

5. Iniciativa global'nogo upravleniya iskusstvennym intellektom // MID KNR. - 2023-10-20 16:05 - https://www.fmprc.gov.cn/rus/wjdt/gb/202310/t20231024_11166700.html (Дата обращения 04.03.2025).

6. German Genshin - Vot chto dumayut 10 liderov mirovogo tekhnologicheskogo rynka ob iskusstvennom intellekte // 15 dekabrya 2023 - <https://vc.ru/u/2830724-andrew-volpert/1004639-samye-vliyatelnye-lyudi-v-mire-ob-ii> (Дата обращения 04.03.2025).

7. Van Czyan'gan, kandidat istoricheskikh nauk - Vliyanie razvitie iskusstvennogo intellekta na processy globalizacii // Teorii i problemy politicheskikh

issledovaniy. 2019. Tom 8. № 1A. S. 193-202.

- <http://www.publishing-vak.ru/file/archive-politology-2019-1/23-wang-jiangang.pdf> (Data obrashcheniya 02.03.2025).

8. E. Yudkovskij - Iskusstvennyj intellekt kak pozitivnyj i negativnyj faktor global'nogo riska //Biologiya, psihologiya, medicina, demografiya i sociologiya E. Yudkovskij — Singularity Institute for Artificial Intelligence Palo Alto, CA Perevod: Aleksej Valer'evich Turchin - <https://spkurdyumov.ru/biology/iskusstvennyj-intellekt/> (Data obrashcheniya 04.03.2025).

9. Dlya chego globalistam iskusstvennyj intellekt? //Ezhenedel'nik \"Zvezda\", 22 yanvarya 2024. - <https://dzen.ru/a/Za41RPqCO1XCuvO3> (Data obrashcheniya 02.03.2025).

10. Iskusstvennyj intellekt i geopolitika. Kak II mozhet povliyat' na vzlyot i padenie nacij? // 8 noyabrya 2023 - RAND Corporation Research <https://doi.org/10.7249/PEA3034-1> 2023 - <https://dzen.ru/a/ZUq-4lmlc2ikf2hq> (Data obrashcheniya 04.03.2025).

11. Ilon Mask: «Samaya razrushitel'naya sila v istorii budet umnee samogoumnogo cheloveka i smozhet vypolnyat' vsyu rabotu», - <https://www.ixbt.com/news/2023/11/06/samaja-razrushitelnaya-sila-v-istorii-budet-umnee-samogo-umnogo-cheloveka-i-smozhet-vypolnjat-vsju-rabotu--ilon-mask.html> (Data obrashcheniya 04.03.2025).

12. Ilon Mask rasskazyvaet o vzryvnom roste iskusstvennogo intellekta i o tom, chto eto znachit dlya budushchego By Brenda Kanana //Obnovleno: 21 iyunya 2024 g. - <https://www.cryptopolitan.com/ru/elon-musk-reveals-ais-explosive-growth/> (Data obrashcheniya 04.03.2025).

13. Upravlenie II v interesah chelovechestva //AI Advisory Body Interim Report: Governing AI for humanity This translation is prepared by MGIMO AI Center. The present work is an unofficial translation for which the publisher accepts full responsibility. Neoficial'nyj perevod dokumenta vypolnen sotrudnikami Centra iskusstvennogo intellekta MGIMO. (Promezhutochnyj otchet). - Dekabr' 2023. - <https://www.un.org/en/ai-advisory-body>. (Data obrashcheniya 05.03.2025).

14. Vasil'ev D.P. Formirovanie mezhdunarodnyh rezhimov upravleniya iskusstvennym intellektom: klyuchevye tendencii i osnovnye aktory. - <https://cyberleninka.ru/article/n/formirovanie-mezhdunarodnyh-rezhimov-upravleniya-iskusstvennym-intellektom-klyuchevye-tendentsii-i-osnovnye-aktory> (Data obrashcheniya 04.03.2025).

Бутко В.Н.,

техника ғылымдарының кандидаты, доцент
marnic08@yandex.kz

*Қостанай әлеуметтік-техникалық университеті
академик З.Алдамжар атындағы, қ.
110000 Қостанай, Қобыланды батыр даңғылы, 27*

ДАМУ ТАРИХЫ ЖӘНЕ ПРОБЛЕМАЛЫҚ МӘНІ ЖАСАНДЫ ИНТЕЛЛЕКТ

***Андамна.** Қазіргі уақытта адамзат жол қиылысында тұр. Генеративті жасанды интеллект (AI) дәуірінде дүниеге келген және барған сайын күшейіп келе жатқан кеңістіктік интеллекттің жаңа технологиясы - AI жүйелерінің үш өлшемді виртуалды және физикалық әлеммен өзара әрекеттесуі - адамзатқа пайда әкеліп қана қоймай, жаңа қауіптер де әкелуі мүмкін. Алгоритмдік бейімділік, жалған ақпарат және автономды жүйелерді пайдалану сияқты жасанды интеллекттің жағымсыз жақтары ерекше назар аударуға тұрарлық. Мұндай проблемаларды тек жаһандық ынтымақтастық арқылы жеңуге болады. Бұл процестің дамуына ұтымды түсіну және оңтайлы әсер ету мүмкіндігі үшін алдымен жасанды интеллекттің (AI) қалыптасу тарихы мен мәнін зерттеу қажет.*

Бұл зерттеудің жоспарланған негізі-жасанды интеллекттің қалыптасу тарихы мен мәнін анықтауға тырысу. Жұмыста зерттеудің келесі ғылыми әдістері қолданылады: аналитикалық, тарихи, логикалық, сараптамалық бағалау.

Жасанды интеллектті қолдану жаһандық ауқымда жылдам қарқынмен кеңейіп, жеделдетілуде. Алайда, бұл технологиялық прогресс жеңіске де, апатқа да айналуы мүмкін, бұл адамзатқа жаһандық сынақ тудырады. Шынында да, жоғары деңгейлі жасанды интеллект өзінің кодын және оңтайландыру мақсатын дербес өзгерте алады. Оның барған сайын интеллектуалды өзін-өзі жетілдіру мүмкіндіктері жаһандық тәуекелдердің айтарлықтай өсуіне әкеледі. Бірақ жасанды интеллект тек сыртқы факторлардың әсерінен дамитын болса да (тек адам инженерлері мен ғалымдарының күшімен), оның интеллектуалды тиімділігі де күрт секіріс жасай алады. Бұл ретте-адамнан толығымен тәуелсіз, өзін-өзі қамтамасыз ететін AI-нің күтпеген "туған" сәті туралы ешкім ешкімге ескертпейді.

Ең қайғылы жағдай – мұның бәрі адамзаттың басым көпшілігінің осы тәуекелдерді жоюға қабілетті интеллектісі әлемді абсолютті басқаруға ұмтылатын заманауи жаһандық сатанистерді - англо-саксондардың "элитасын" және ең алдымен жаңа Британ империясының басқарушы "элитасын" шектеуге, "нөлге келтіруге" тырысатын жағдайларда орын алады. Олар терроризмді, соғыстарды, неонацизмді, ЛГБТ-ны насихаттау, білім беруде мақал-мәтелдер жүйесін енгізу және т. б. арқылы барлық деңгейдегі мәдениетті, білім мен ғылымды бұзады.

Біз әлі күнге дейін жетілдірілген AI жүйелерінің өзін қалай ұстайтынын толық түсінбейміз. Сондықтан сарапшылардың көпшілігі әлемдік қауымдастыққа жасанды интеллектті тиімді басқару қажет деп санайды. Алайда, бұл жерде жаңа проблема туындайды-қандай әлемдік күштер бірінші болып жетістікке жетеді (ал екіншісі болмайды!) жасанды интеллектті – қазір әлемде үстемдік ететін зұлымдық күштерін немесе қазір зұлымдық күштерінен "Жер әлемін басқару рулін" ұстап алуға тырысып жатқан

жақсылық күштерін бақылауға алу керек. Сондықтан, ең алдымен, жаһандық басқарудың факторы ретінде АИ-ны одан әрі мұқият зерттеу қажет.

Түйін сөздер: *АИ тарихы, мәні, жағымсыз жақтары, жаһандық ынтымақтастық, АИ коды мен мақсаты, АИ өзін-өзі жетілдіру, жаһандық тәуекелдер, өзін-өзі қамтамасыз ету, АИ басқару, бақылау.*

Butko V. N.,

Candidate of Technical Sciences, Associate Professor,
marnic08@yandex.kz

*Kostanay Social-Technical University named after
academician Z. Aldamzhar,
110000 Kostanay, Kobylandy Batyr Ave. Kostanay, 27*

HISTORY OF FORMATION AND PROBLEM ESSENCE ARTIFICIAL INTELLIGENCE

Abstract. *Nowadays Mankind is at a crossroads. Born in the era of generative artificial intelligence (AI) and gaining more and more strength, the new technology of spatial intelligence - the interaction of AI systems with the three-dimensional virtual and physical world - can bring humanity not only benefits, but also new threats. The negative aspects of AI, such as algorithmic bias, misinformation, and the use of autonomous systems, deserve special attention. Such challenges can only be overcome through global cooperation. And for the possibility of rational understanding and optimal influence on the development of this process, it seems necessary to first explore the history of formation and the essence of artificial intelligence (AI).*

The planned basis of this research is to try to find out the history of formation and the essence of artificial intelligence. The following scientific methods of research are used in this work: analytical, historical, logical, expert evaluations.

The application of artificial intelligence is expanding and accelerating at a rapid pace on a global scale. However, this technological progress can turn out to be both a victory and a disaster, forming a global challenge to humanity. Indeed, top-level AI is capable of self-modifying its code and optimization purpose. And its capacity for ever-faster intellectual self-improvement is causing global risks to increase significantly. But even if AI will develop only due to external factors (solely by human engineers and scientists), its intellectual efficiency can also make a sharp leap. At the same time, no one will warn anyone about the unexpected moment of the “birth” of self-sufficient AI, already completely independent of Man.

And the most tragic thing is that all this is happening in conditions when the intelligence of the overwhelming part of Humanity, capable of eliminating these risks, is very diligently trying to limit, “nullify” the modern Global Satanists, who strive for

absolute control of the world - the 'elite' of the Anglo-Saxons and, above all, the ruling "elite" of the New British Empire. They destroy culture, education and science at all levels by unleashing terrorism, wars, propaganda of neo-Nazism, LGBT, introduction of the notorious Bologna system in education, etc. etc. etc.

We still do not fully understand how advanced AI systems behave. That is why most experts agree that the world community needs to effectively manage AI. However, a new problem arises here - which world forces will be the first (and there will be no second!) to put artificial intelligence under their control - the Forces of Evil, which now prevail in the world, or the Forces of Good, which are now only trying to take over the "Rudder of the Earth's World" from the Forces of Evil. Therefore, further thorough research of AI as a factor, first of all, of Global Governance is necessary.

Keywords: *AI history, essence, negative sides, global cooperation, AI code and purpose, AI self-improvement, global risks, self-sufficiency, AI management, control.*