

К истории искусственного интеллекта

Ю.Е. Поляк

*Центральный экономико-математический институт
Российской академии наук, Москва*

Аннотация. Излагается краткая история возникновения и развития искусственного интеллекта (ИИ), включая Дартмутскую конференцию 1956 года, на которой были заложены основы современного ИИ, и где был впервые использован этот термин. Кроме того, упомянуты мифы и сказания, где отражались мечты о рукотворных помощниках, наделенных искусственным разумом. Описаны работы предшественников современного ИИ – философов и изобретателей. Уделено внимание общественным инициативам («Римский призыв к этике ИИ», открытое письмо Pause Giant AI Experiments), а также нормативным актам, регулирующим ИИ в странах всех континентов.

Ключевые слова: искусственный интеллект; история; машина Бэббиджа; тест Тьюринга; экспертные системы; нейронные сети

On the history of artificial intelligence

Y.E. Polak

Central Economics and Mathematics Institute. Moscow, Russia

Abstract. A brief history of the emergence and development of artificial intelligence (AI) is presented, including the Dartmouth Conference of 1956, where the foundations of modern AI were laid and where this term was first used. In addition, myths and legends are mentioned, reflecting dreams of man-made assistants endowed with artificial intelligence. The works of the predecessors of modern AI - philosophers and inventors - are described. Attention is paid to public initiatives (the Rome Call for AI Ethics, the open letter Pause Giant AI Experiments), as well as regulations governing AI in countries on all continents.

Keywords: artificial intelligence; history; babbage machine; turing test; expert systems; neural networks

Искусственный интеллект (ИИ) как наука находится на стыке математики, кибернетики, биологии, психологии. Он изучает технологии, которые позволяют человеку писать «интеллектуальные» программы и учить компьютеры решать задачи самостоятельно. Главная задача ИИ —

понять, как устроен человеческий интеллект, и смоделировать его. Среди подразделов ИИ – робототехника, наука о компьютерном зрении, обработка естественного языка, машинное обучение.

Началом искусственного интеллекта многие считают конференцию 1956 года в Дартмуте, на которой были заложены основы современного ИИ, и где был впервые использован этот термин. Его предложил Джон Маккарти, который наряду с Марвином Мински, Клодом Шенноном и Натаниэлем Rochesterом считается отцом-основателем ИИ. Дартмутская конференция стала отправной точкой для многих исследований и разработок в области ИИ.

Но еще задолго до 1956 года философы пытались понять природу человеческого разума. Аристотель разработал формальную логику, которая позже стала основой для алгоритмов. Герон Александрийский создал механические игрушки, такие как театр теней и самодвижущаяся тележка, открывающиеся ворота. Появлялись и другие устройства, проявляющие «разум». В мифах Древней Греции встречается автоматон — это кукла, выполняющая действия по заданному алгоритму. Другой пример – Пандора, девушка, созданная Гефестом из земли и воды по приказу Зевса. Искусственные создания фигурируют в мифах о Кадме и о Пигмалионе. В культуре еврейского народа встречаются мистические големы, созданные из неживой материи. Они выступают прообразом современного искусственного интеллекта.

Прообразы ИИ часто встречается и в русских сказках: герои пользуются различными формами «искусственного интеллекта», чтобы кто-то сделал за них невыполнимую работу. Это управляемые голосом горшочек каши, которую нужно сварить, скатерть-самобранка, из которой появлялось всё необходимое, изба Бабы Яги, сани из сказки «По щучьему велению». В той же сказке печь играет роль роботизированного такси.

В XVII веке философы заговорили о механизации мышления, о возможности вдохнуть жизнь в неживые предметы. Рене Декарт предложил механистическую модель человека, сравнивая его с часами. Готфрид Лейбниц считал, что мысли человека можно изобразить при помощи символов и схем. Идеи об искусственных людях и мыслящих машинах стали популярной темой в художественной литературе. Среди таких персонажей – человекоподобное существо гомункулус из трагедии Гёте «Фауст», выращенное в колбе путём химических реакций, Франкенштейн Мэри Шелли. Идеи искусственной жизни нашли отражение в произведениях «Шахматист Мельцеля» Э.По, «Дарвин среди машин» С.Батлера, «R.U.R.» К.Чапека. Из этой пьесы в обиход вошло слово «робот», предложенное братом автора Й.Чапеком.

Реальность старалась поспевать за фантастикой. Английский математик Ч.Бэббидж в 1834 году начал работу над универсальной вычислительной машиной, которую учёный назвал аналитической. Её

справедливо считают прообразом современного цифрового компьютера. Несмотря на то, что Бэббидж подробно описал конструкцию аналитической машины и принципы её работы, она так и не была построена при его жизни из-за недостатка средств. Работу закончил сын изобретателя. В 1906 году устройство показало свою работоспособность.

Знакомство с Бэббиджем вдохновило леди Августу Аду Лавлейс написать программу подсчёта чисел Бернулли для аналитической машины, еще не реализованной. В 1843 году она создала описание ЦВМ и инструкции по программированию к ней. Лавлейс считается первым программистом в истории. В ее честь назван язык программирования Ада.

В 1950 г. А.Тьюринг предложил тест для определения, сможет ли компьютерная программа имитировать человеческий интеллект. Человек ведет текстовый диалог с машиной и другим человеком, не зная, где компьютер. Если после серии вопросов и ответов невозможно определить, с кем он общался, то считается, что робот прошел тест. Тест Тьюринга стал одним из методов оценки искусственного интеллекта. Идея опубликована в статье «Вычислительные машины и разум» [1], см. также русское издание [2].

В июне 2024 года тест Тьюринга прошёл чат-бот с генеративным искусственным интеллектом ChatGPT-4 от OpenAI. 54% участников решили, что беседуют с реальным человеком. Полвека назад некоторым казалось, что первой программой, которая прошла тест Тьюринга, была «Элиза» (ELIZA) — первый в истории цифровой чат-бот, программа виртуального собеседника, разработанная Д.Вейценбаумом в 1966 году. Она названа в честь героини пьесы Шоу «Пигмалион».

Алгоритм работы заключался в формировании ответов на основе реплик человека, вводимых с телетайпа. Программа выделяла значимые слова и подставляла их в шаблон ответа, создавая имитацию разговора с психотерапевтом, например:

User: *I'm depressed much of the time.*

ELIZA: *I am sorry to hear you are depressed.*

User: *It's true. I am unhappy.*

ELIZA: *Can you explain what made you unhappy?*

При затруднениях выдавались ответы вроде «Я понимаю» или «Продолжайте». Иллюзия, что «Элиза» понимает собеседника, иногда была настолько правдоподобной, что программе приписывали понимание, сочувствие и другие человеческие качества.

В 1970-е годы Вейценбаум писал статьи и книги, предупреждая об опасностях искусственного интеллекта. Учёный полагал, что ИИ служит «индексом безумия нашего мира», а передавал многие решения

компьютерам, люди создали менее рациональный мир. В книге [3] он призывал ограничить машины задачами, где нужны только вычисления.

Опасения Вейценбаума не напрасны. Даже мощнейшим нейросетям не следует давать право принимать важные решения, ведь они лишены представлений об этике. Беспилотные автомобили могут попасть в ситуацию, угрожающую человеческим жизням. Вместо посещения терапевта порой направляют к чат-ботам. В США ИИ используют для оценки рисков при решениях о приговоре, освобождении под залог, не говоря уже о предоставлении кредита. Регулирование использования нейросетей только начинается.

Вернемся к истории ИИ. В 1960-е годы были также созданы первые алгоритмы для таких задач как распознавание образов, обработка естественного языка. Тогда же начались исследования в области нейронных сетей, которые стали основой современных методов машинного обучения. Начались разработки моделей, имитирующих работу человеческого мозга и обучаемых на основе данных.

Проект MYCIN (MYcological Consultant for Infectious Diseases) разрабатывался в 1964-1974 годах в Стэнфордском университете как одна из первых экспертных систем для диагностической поддержки при инфекционных заболеваниях. Система оперировала с помощью базы знаний из 600 правил. Программа задавала врачу ряд текстовых вопросов. В результате система предоставляла список подозреваемых бактерий и рекомендовала курс лечения. По данным Stanford Medical School, MYCIN предлагает приемлемую терапию в 69% случаев, что лучше, чем у профессиональных экспертов. Это вызвало новое обсуждение важных этических вопросов: какова ответственность врачей, если система ставит неверный диагноз; как обеспечить конфиденциальность и безопасность данных, и т.д.

В 1997 году шахматный компьютер Deep Blue от IBM победил чемпиона мира Г.К.Каспарова. Это стало важным моментом в истории ИИ. В 2015-2017 годах программа AlphaGo для игры в го, разработанная компанией DeepMind, неоднократно побеждала в матчах сильнейших игроков мира Фань Хуэя, Ли Седола, Кэ Цзе.

Активно разрабатывал проблемы искусственного интеллекта Яндекс, одна из ведущих ИТ-компаний [5]. Его американскую дочернюю фирму Avride (ранее Yandex Self Driving Group), которая занималась разработкой программного и аппаратного обеспечения для беспилотных автомобилей (испытания таких автомобилей Яндекс провёл ещё в 2017 г.), аналитики Morgan Stanley называли одним из лидеров на мировом рынке и оценивали в \$7 млрд. Шестиколёсные роботы-курьеры Яндекс.Ровер работают с 2019 г. (прототип был создан в 2015 г.). Они доставляют еду, посылки, почту и оказались особенно полезны в период самоизоляции. В 2021-22 гг. дроны Яндекса развозили еду по кампусам университетов Огайо и

Аризоны¹. Другая разработка Яндекса – виртуальный ИИ-ассистент Алиса, появившаяся в 2017 г. как голосовой помощник. С 2018 года в продажу выходит линейка «умных колонок», беспроводных bluetooth-устройств, а с 2022 г. – умный телевизор с той же Алисой для голосового управления. В 2023 г. Алиса научилась писать тексты, предлагать идеи и составлять планы. Наконец, с мая 2025 года Алиса встроена в поисковую машину Яндекса, что даёт пользователю новое качество поиска: в ответ на запрос он получает небольшой структурированный текст с иллюстрациями.

Сейчас ежемесячная аудитория Алисы — более 70 млн человек, половина населения страны.

Давно возник интерес к проблеме взаимоотношений человека с машиной. Еще в 1942 г. в рассказе «Хоровод» А.Азимов сформулировал «Три закона робототехники», по мотивам категорического императива И.Канта.

1. Робот не может причинить вред человеку или своим бездействием допустить, чтобы человеку был причинён вред.

2. Робот должен повиноваться всем приказам, которые даёт человек, кроме тех случаев, когда эти приказы противоречат Первому Закону.

3. Робот должен заботиться о своей безопасности в той мере, в которой это не противоречит Первому или Второму Законам.

Впоследствии, в романе «Роботы и Империя» (1985), Азимов добавил еще один закон.

0. Робот не может нанести вред человечеству или своим бездействием допустить, чтобы человечеству был нанесён вред.

Обратим внимание на книгу М.Пасквинелли [4], где с марксистских позиций исследуется ИИ как не только научно-технический, но в первую очередь социально-экономический феномен. Ссылаясь на работы Ч.Бэббиджа и А.Тьюринга, автор указывает, что «внутренний код ИИ формируется не имитацией биологического интеллекта, а интеллектом труда и социальных отношений», тайна и чудо искусственного интеллекта — это всего лишь автоматизация труда в высшей степени, а не интеллект как таковой, а идея о том, что ИИ может однажды стать автономным или разумным - чистая фантазия. По его словам, в начале ХХI века беспилотные транспортные средства и другие артефакты побудили иначе взглянуть на ручные навыки и обнаружить, что самое ценное в труде никогда не сводилось к ручному компоненту, а включало когнитивные и кооперативные аспекты². То есть нужно признать, что благодаря

¹ <https://www.ixbt.com/news/2021/08/19/roboty-jandeksa-nachali-kormit-studentov-v-ssha.html>

² Вспоминается. Замечательный программист и педагог Александр Львович Брудно утверждал, что в любом труде гораздо больше ценится интеллектуальная составляющая и приводил «математическое доказательство». $200 \text{ (часов)} * 1/20 \text{ (л.с.)} * 0.75 \text{ (квт)} * 4 \text{ (коп)} = 30$. Здесь 4 копейки – стоимость киловатт-часа для массового потребителя в СССР. Месячные зарплаты трудящихся в сотни раз превышали 30 копеек, которые могла бы стоить их физическая энергия.

исследованиям ИИ водители грузовиков примкнули к интеллектуалам. ИИ навязывает новое разделение труда, в котором коллективная деятельность людей оказывается пролетаризованной, а произведенные знания узурпируются небольшим числом игроков, находящихся на вершине цифровой экономики.

Пасквинелли пишет, что искусственный интеллект продолжает «кодировать социальные иерархии и дискриминировать рабочую силу». Он утверждает, что сущность искусственного интеллекта не в том, чтобы воспроизводить человеческий разум, а в новых ступенях автоматизации умственного труда и в инструментах контроля («хозяйский глаз» в названии книги отсылает к выражению Ф.Энгельса, которым тот описывал контроль над рабочими во времена промышленной революции). По Пасквинелли, нет причин бояться, что ИИ выйдет из-под контроля, обретет автономность и уничтожит человечество — это противоречит его природе и заложенным в него прагматичным и социально обусловленным задачам.

У католической церкви есть документ, официально обозначающий на высшем уровне её позицию по ИИ. 28 февраля 2020 года Папская академия защиты жизни вместе с Microsoft и IBM подписали The Rome Call for AI Ethics - «Римский призыв к этике ИИ»³, к которому присоединились мусульмане-сунниты и иудеи. Он определяет пункты, которым должен следовать искусственный интеллект. Это подотчетность, прозрачность, отсутствие дискриминации, беспристрастность, надежность, конфиденциальность и безопасность. На конференции «Этика ИИ во имя мира» в Хиросиме в июле 2024 г. к «Римскому призыву» присоединились 11 мировых религий (буддизм, индуизм, зороастризм, бахаизм и др.)⁴.

Вот еще одна общественная инициатива. В начале февраля 2023 г. чат-бот ChatGPT стал самым быстрорастущим сервисом в истории. Всего через два месяца после запуска им воспользовались около 100 млн активных пользователей. Такое бурное развитие вызвало беспокойство ряда специалистов и бизнесменов, связанных с ИИ. 22 марта 2023 организация Future of Life опубликовала открытое письмо Pause Giant AI Experiments⁵ с призывом остановить исследование мощных систем ИИ. Среди сотен подписей (на май 2025 г. их свыше 33 тысяч) выделяются имена И.Маска и С.Возняка. В письме обращается внимание, что по результатам крупных исследований, системы ИИ с интеллектом, сопоставимым с человеческим, могут представлять опасность для общества. Авторы текста призывают все лаборатории, исследующие искусственный интеллект, немедленно приостановить разработку и обучение нейросетей, пока не появятся общие протоколы безопасности. «Эта пауза должна быть публичной и действительной. Если же подобная

³ <https://www.romecall.org/the-call/>

⁴ <https://www.romecall.org/ai-ethics-for-peace-hiroshima-july-10th-2024/>

⁵ <https://futureoflife.org/open-letter/pause-giant-ai-experiments/>

приостановка не может быть сделана быстро, правительства государств должны вмешаться и ввести мораторий», подчеркивается в тексте письма⁶.

Итак, за последние годы ИИ превратился в революционную технологию с бесчисленными применениями в реальном мире. На передний план вышли такие темы как управление моделями, безопасность и риски. Правительства по всему миру отреагировали на это предложениями по регулированию ИИ. Более 70 стран приняли свыше 1000 политических инициатив и правовых основ для решения проблем, связанных с безопасностью и управлением ИИ⁷.

Одной из первых стран, разработавших национальную стратегию в области ИИ, стал Китай. В июле 2017 года опубликован План развития искусственного интеллекта нового поколения⁸. Этот план ставил амбициозные цели в области развития ИИ, рассчитанные до 2030 года.

В мае 2024 года предложен проект закона об искусственном интеллекте Китайской Народной Республики⁹, который накладывает юридические обязательства на разработчиков и пользователей. В сентябре 2024 года опубликована AI Safety Governance Framework¹⁰ – схема управления безопасностью ИИ, в которой представлены рекомендации по этичному и безопасному развитию технологий ИИ.

В 2022 году Япония опубликовала Национальную стратегию в области ИИ (National AI Strategy)¹¹, согласно которой правительство в рамках концепции «гибкого управления» возлагает на частный сектор добровольные усилия по саморегулированию. 16 февраля 2024 г. опубликован проект «Основного закона о содействии ответственному использованию ИИ»¹². Выпущено руководство по ИИ для бизнеса¹³.

В Австралии в 2024 г. опубликован добровольный стандарт безопасности ИИ, и начала действовать «политика ответственного использования ИИ в правительстве»¹⁴, направленная на то, чтобы представить правительство примером безопасного и ответственного использования технологий ИИ.

В Индии была создана рабочая группа для выработки рекомендаций по этическим, юридическим и социальным вопросам, связанным с ИИ. Согласно Национальной стратегии развития ИИ в стране (National Strategy

⁶ <https://www.rbc.ru/life/news/6424457c9a7947ebee7f7534>

⁷ <https://www.mindfoundry.ai/blog/ai-regulations-around-the-world>

⁸ <https://digichina.stanford.edu/work/full-translation-chinas-new-generation-artificial-intelligence-development-plan-2017>

⁹ <https://cset.georgetown.edu/publication/china-ai-law-draft>

¹⁰ https://english.www.gov.cn/news/202409/10/content_WS66df9f30c6d0868f4e8eac91.html

¹¹ <https://www8.cao.go.jp/cstp/ai/aistrategy2022en.pdf>

¹² <https://www.dlapiper.com/en-gb/insights/publications/2024/10/understanding-ai-regulations-in-japan-current-status-and-future-prospects>

¹³ https://www.meti.go.jp/shingikai/mono_info_service/ai_shakai_jisso/pdf/20240419_9.pdf

¹⁴ <https://www.cglaw.com.au/government-sets-policy-for-its-use-of-ai>

for AI)¹⁵, Индия надеется стать «гаражом ИИ» для стран с формирующимся рынком и развивающихся стран.

В январе 2024 года президент ОАЭ объявил о создании Совета по искусственному интеллекту и передовым технологиям (Artificial Intelligence and Advanced Technology Council, AIATC)¹⁶.

В Саудовской Аравии в 2023 году были опубликованы «Принципы этики ИИ» (AI Ethics Principles September 2023)¹⁷.

17 декабря 2023 года Израиль представил свою «всеобъемлющую политику в области регулирования и этики искусственного интеллекта»¹⁸.

22 июня 2022 года правительство Канады объявило о запуске второго этапа Всеканадской стратегии в области искусственного интеллекта¹⁹, на который выделено более 443 миллионов долларов. Он направлен «на привлечение талантов мирового уровня и передовых исследовательских мощностей для коммерциализации и внедрения канадских идей и знаний».

В Бразилии юридическая основа для развития искусственного интеллекта (Marco Legal da Inteligência Artificial)²⁰ появилась 10 декабря 2024 года, когда Федеральный сенат одобрил закон №2338/23.

В США отдельные штаты имеют собственное законодательство. Так, закон штата Калифорния о прозрачности ИИ (California AI Transparency Act SB 942)²¹, начинающий действовать 1 января 2026 года, обязывает компании раскрывать информацию об использовании генеративных ИИ-систем. В 2025 году 45 штатов внесли более 550 законопроектов, связанных с ИИ, в том числе продолжают попытки принять закон об ответственности разработчиков ИИ за причинённый вред. Но в мае республиканская партия предложила законопроект, запрещающий штатам и другим местным органам власти регулировать «модели искусственного интеллекта, системы искусственного интеллекта или автоматизированные системы принятия решений в течение 10 лет с даты принятия закона»²².

30 октября 2023 г. президент США Джозеф Байден подписал указ №14110 о безопасной, надёжной и заслуживающей доверия разработке и использовании искусственного интеллекта (The Executive Order #14110 on Safe, Secure, and Trustworthy Artificial Intelligence)²³. Указ, направленный

¹⁵ <https://niti.gov.in/sites/default/files/2019-01/NationalStrategy-for-AI-Discussion-Paper.pdf>

¹⁶ <https://www.mediaoffice.abudhabi/en/government-affairs/in-his-capacity-as-ruler-of-abu-dhabi-uae-president-issues-law-establishing-artificial-intelligence-and-advanced-technology-council>

¹⁷ <https://sdaia.gov.sa/en/SDAIA/about/Documents/ai-principles.pdf>

¹⁸ https://www.gov.il/en/pages/ai_2023

¹⁹ <https://www.canada.ca/en/innovation-science-economic-development/news/2022/06/government-of-canada-launches-second-phase-of-the-pan-canadian-artificial-intelligence-strategy.html>

²⁰ <https://www.clarkemodet.com/en/legislative-news/brasil-aprueba-un-nuevo-marco-legal-para-regular-la-ia/>

²¹ <https://digitalpolicyalert.org/event/18058-announced-bill-on-consumer-protection-in-generative-artificial-intelligence-sb-942>

²² <https://rossaprimavera.ru/news/0e6f8ad8>

²³ <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>

на регулирование технологий искусственного интеллекта обязывал разработчиков ИИ-алгоритмов предоставлять правительству данные, полученные в ходе тестирования своих продуктов на безопасность и защиту их технологий²⁴. Республиканская партия раскритиковала указ, заявив, что он тормозит развитие инноваций и накладывает чрезмерные ограничения на бизнес.

Президент Дональд Трамп в первый день своей каденции 20 января 2025 года одной подписью отменил сразу 78 указов и распоряжений Байдена, включая указ 14110²⁵. Отмена указа сигнализирует о более либеральном подходе администрации Трампа к регулированию ИИ²⁶. Новый президент сосредоточится на стимулировании инноваций и ослаблении регулирования, в отличие от политики усиленного контроля, проводимой Байденом. И 7 апреля на сайте Белого дома появился информационный бюллетень «Устранение барьеров для использования и закупок искусственного интеллекта на федеральном уровне» (Eliminating Barriers for Federal Artificial Intelligence Use and Procurement)²⁷. Теперь правительство больше не будет вводить ненужные бюрократические ограничения на использование искусственного интеллекта в исполнительной власти.

13 марта 2024 г. Европейский парламент одобрил комплексный закон по регулированию искусственного интеллекта²⁸. «Целью закона является предоставление разработчикам и специалистам по внедрению ИИ четких требований и обязательств в отношении конкретных видов использования этих технологий», - указывается в заявлении Еврокомиссии. Создается правовая база, которая устраняет риски, связанные с искусственным интеллектом, и позволяет Европе играть ведущую роль в данной области.

В России с 2019 года действует «национальная стратегия развития искусственного интеллекта на период до 2030 года», введенная указом от 10.10.2019 № 490²⁹. В ней говорится об основных принципах развития и использования технологий ИИ, целях такой работы, приоритетных направлениях, поддержке научных исследований, подготовке кадров, регулировании общественных отношений в этой сфере.

Вторая версия «стратегии» появилась в 2024 г. (указ от 15.02.2024 №124)³⁰. В документе приводится ряд целевых показателей. Так, ежегодный объем оказанных услуг по разработке и реализации решений в

²⁴ <https://3dnews.ru/1095218/prezident-ssha-predstavil-ukaz-v-sfere-regulirovaniya-ii>

²⁵ <https://kod.ru/trump-deny-law-secure-ai-slow>

²⁶ <https://www.reuters.com/technology/artificial-intelligence/trump-revokes-biden-executive-order-addressing-ai-risks-2025-01-21>

²⁷ <https://www.whitehouse.gov/fact-sheets/2025/04/fact-sheet-eliminating-barriers-for-federal-artificial-intelligence-use-and-procurement>

²⁸ <https://ria.ru/20240313/zakon-1932729659.html>

²⁹ <http://publication.pravo.gov.ru/Document/View/0001201910110003>

³⁰ <http://publication.pravo.gov.ru/document/0001202402150063>

области ИИ к 2030 году должен вырасти как минимум до 60 млрд рублей (в 2022 году – 12 млрд). Количество выпускников вузов с образованием в сфере искусственного интеллекта планируется увеличить за тот же период с 3 до 15.5 тыс. человек в год. Есть и более трудноформализуемые цели. Уровень доверия граждан к технологиям искусственного интеллекта в 2030 году должен вырасти не менее чем до 80% по сравнению с 55% в 2022 году, а доля приоритетных отраслей экономики с высокой готовностью к внедрению ИИ – с 12% до 95%. Что ж, через пять лет можно будет сравнить пожелания с реальностью.

Литература

1. Turing A. Computing machinery and intelligence // Mind: Oxford University Press, 1950. — No. 59. — P. 433-460.
2. Тьюринг А. Может ли машина мыслить? (С приложением статьи Дж. фон Неймана «Общая и логическая теория автоматов»). — М.: Государственное издательство физико-математической литературы, 1960. - 67 с.
3. Weizenbaum J. Computer power and human reason: from judgment to calculation. San Francisco, 1976 - 300 p.
4. Pasquinelli M. The Eye of the Master: A Social History of Artificial Intelligence. - Verso, 2023 - 272 p.
5. Поляк Ю.Е. Пять пятилеток и один год // Научный сервис в сети Интернет. – 2023. – № 25. – С. 308-319. – DOI 10.20948/abrau-2023-17.