

RU

«Реальная этика» искусственного интеллекта и проблема морального эскапизма

Дедюлина М. А.

Аннотация. Цель данного исследования – построение концепции «реальной этики искусственного интеллекта (ИИ)», обеспечивающей переход от доминирующей парадигмы проектирования (с приоритетом вытесняющих отношений) к расширению фронезиса на основе эксплицированного принципа У. Черчилля и диагностированных форм морального эскапизма. Научная новизна заключается в том, что автор обращается к принципу Черчилля впервые применительно к ИИ; раскрыт моральный эскапизм в практике использования ИИ; построена типология опосредования фронезиса; доказан приоритет вытесняющих отношений в доминирующей парадигме проектирования ИИ; сформулированы принципы «реальной этики ИИ» для расширения фронезиса. В результате проведенного исследования обоснована применимость принципа Черчилля к анализу ИИ; выявлены основные формы морального эскапизма и его проявления в практике использования ИИ; построена и обоснована типология технологического опосредования фронезиса; показано, что доминирующая парадигма проектирования ИИ систематически отдаёт приоритет вытесняющим и ослабляющим отношениям; сформулированы принципы «реальной этики ИИ», ориентированные на расширение, а не замещение практического разума человека.

EN

“Real ethics” of artificial intelligence and the problem of moral escapism

M. A. Dedyulina

Abstract. The aim of the study is to construct a concept of “real ethics of artificial intelligence (AI)” that facilitates a shift away from the dominant design paradigm – characterized by privileging displacing relations – toward the expansion of phronesis on the basis of an explicated Churchillian principle and a diagnosis of forms of moral escapism. The scientific novelty lies in the first application of the Churchillian principle to AI; the identification of moral escapism in AI practice; the development of a typology of mediation of phronesis; the demonstration that the prevailing AI-design paradigm systematically prioritizes displacing and attenuating relations; and the formulation of principles of “real ethics of AI” aimed at expanding, rather than replacing, human practical reason. The study shows that W. Churchill’s principle is applicable to AI analysis; it identifies the main forms of moral escapism and their manifestations in AI use. The work constructs and substantiates a typology of technological mediation of phronesis, and demonstrates that the dominant design paradigm favors relations that displace or weaken human practical agency. Finally, the paper formulates principles of “real ethics of AI” oriented toward augmentation of human practical reason rather than its substitution.

Введение

Стремительное проникновение генеративных нейросетей, рекомендательных алгоритмов и систем поддержки принятия решений в повседневность ставит современную этику перед фундаментальным вопросом: как нам следует жить с технологиями, которые всё чаще принимают решения за нас? В публичном дискурсе сталкиваются два противоположных нарратива. Первый гласит, что искусственный интеллект поможет нам принимать лучшие моральные решения, устраняя человеческую предвзятость и обеспечивая беспристрастный расчёт последствий. Второй, напротив, выражает тревогу, что ИИ лишает нас моральной самостоятельности, превращая в пассивных потребителей машинных суждений (Vallor, 2015, p. 107-124).

Актуальность настоящего исследования обусловлена тем, что современные системы ИИ проектируются преимущественно как «устройства», сводящие практическую мудрость к потреблению, что создаёт риск систематической моральной дисквалификации – утраты способности к самостоятельному этическому суждению (Schemmer, Kühl, Satzger, 2021, p. 108; Vallor, 2017, p. 163). В связи с этим возникает необходимость

в анализе не только абстрактных принципов, но и материальных, технологических условий, в которых совершаются моральные действия (Borgmann, 2006, p. 11). Проблема исследования заключается в противоречии между требованием ответственного, укоренённого в материальной среде морального действия и устойчивым искушением переложить ответственность за моральный выбор на внешние инстанции – экспертов, общественное мнение или алгоритмы (Hertzberg, 2002, p. 257). В эпоху ИИ это противоречие обретает новую, технологически усиленную форму.

Для достижения цели исследования необходимо решить несколько задач: 1) эксплицировать принцип Черчилля в интерпретации Альберта Боргманна и показать его применимость к анализу ИИ-технологий; 2) раскрыть содержание понятия «моральный эскапизм» по Ларсу Гертцбергу и выявить его проявления в практике использования ИИ; 3) представить и подробно обосновать типологию технологического опосредования фронезиса; 4) показать, что доминирующая парадигма проектирования ИИ отдаёт приоритет вытесняющим и ослабляющим отношениям; 5) сформулировать принципы «реальной этики ИИ», направленной на расширение фронезиса.

В работе используются философско-аналитический метод, включающий концептуальный анализ и экспликацию понятий, постфеноменологический подход, анализирующий технологическое опосредование человеческого опыта с акцентом на мультистабильность технологий (Ihde, 1979; Verbeek, 2005; Rosenberger, 2014), кейс-метод, предполагающий разбор конкретных примеров из современной технологической практики, и герменевтический анализ ключевых текстов.

Теоретическая база исследования включает работы Альберта Боргманна (Borgmann, 1992; 2006), Ларса Гертцберга (Hertzberg, 2002), Эндрю Зельни (Zelny, 2025), Шеннон Валлор (Vallor, 2015), а также постфеноменологические исследования Дона Иде (Ihde, 1979; 1999), Питера-Пола Вербика (Verbeek, 2005; 2011), а также Роберта Розенбергера (Rosenberger, 2014), в том числе в соавторстве с Вербиком (Rosenberger, Verbeek, 2015).

Так, в работах американского философа Альберта Боргманна (Borgmann, 1992; 2006) критикуется абстрактность основной англо-американской философии, которая, по его мнению, слишком оторвана от повседневной жизни. Он предлагает философию вовлечённую, а не отстранённую, и более синтетическую, чем аналитическую (Borgmann, 2006, p. 2-3). Центральным понятием Боргманна в работе «Истинная американская этика: взять на себя ответственность за нашу страну» («Real American Ethics: Taking Responsibility for Our Country») является «принцип Черчилля», основанный на высказывании Уинстона Черчилля в 1943 году при обсуждении восстановления здания Палаты общин после бомбардировок: «Мы строим наши здания, а затем наши здания формируют нас (здесь и далее перевод наш. – М. Д.)» (Borgmann, 2006, p. 5). В интерпретации Боргманна (Borgmann, 1992, p. 295) этот принцип означает, что созданные человеком материальные и социальные структуры – физическая среда, институты, системы, инфраструктура – не являются нейтральным фоном, а наполнены моральным содержанием и влияют на то, кто мы есть и как мы живём. А. Боргманн (Borgmann, 1992) полагает, что нельзя отождествлять понятия «вещь» и «устройство». Вещь требует практики, внимания, умелого и активного участия человека. Для подтверждения этой мысли он описывает работу музыкального инструмента. Чтобы извлечь из него музыку, необходимо активно взаимодействовать с ним и применять свои знания и навыки. Как пишет Боргманн, «вещь, в том смысле, в котором я хочу использовать этот термин, обладает понятным и доступным характером и требует умелого и активного участия человека. Вещь требует практики, в то время как устройство приглашает к потреблению» (Borgmann, 1992, p. 296). Устройство, напротив, сводит практику к простому товару потребления. Примером может служить стереосистема: простым щелчком мыши музыка становится доступной слушателю без какого-либо дополнительного участия, кроме пассивного прослушивания. Философ подчёркивает: «Мы не чувствуем себя обязанными присутствию стереосистемы так, как чувствуем себя, слушая музыканта. Мы уважаем музыканта, мы владеем стереосистемой» (Borgmann, 1992, p. 295). Поэтому вещи, по Боргманну (Borgmann, 1992, p. 296-297), составляют «господствующую реальность» – реальность, которая требует от нас ответа, внимания, уважения. Устройства же создают «одноразовую реальность» – реальность, в которой практика сводится к потреблению, а результат становится товаром, лишённым внутренней ценности. Однако он предупреждает, что «одноразовая реальность сводит практику к простому товару: конечная цель – лишь то, что человек получает в итоге, практически не заботясь о процессе, который в неё вкладывается» (Borgmann, 1992, p. 298). Он признаёт риск любого материалистического этического проекта: «Один из рисков любого материалистического или институционалистского этического проекта заключается в том, что он переложит ответственность с отдельного человека на монолитную структуру» (Borgmann, 2006, p. 120). Отдалённые, непроницаемые институты, перед которыми простые люди чувствуют себя бессильными, могут порождать моральный эскапизм. Поэтому философ предлагает «экономику отдельного домохозяйства» – практику осознанного управления своей ближайшей средой, начиная с того, куда поставлен телевизор (Borgmann, 2006, p. 120-121).

Идеи Боргманна можно объединить с идеями финского философа Ларса Гертцберга, который под влиянием позднего Л. Витгенштейна, в своей работе «Моральный эскапизм и прикладная этика» («Moral Escapism and Applied Ethics») вводит понятие «морального эскапизма», под которым понимается искушение избежать окончательной ответственности за моральный выбор путём обращения к внешним инстанциям: общественному мнению, моральному консенсусу или экспертизе (Hertzberg, 2002, p. 252). Гертцберг пишет: «Моральная серьёзность, по крайней мере в одном смысле этого слова, как раз и характеризуется признанием того, что существуют некоторые вопросы, в отношении которых я не могу переложить окончательную ответственность на кого-либо другого. Это выражается в наших рассуждениях о таких вопросах, как вопросы совести. Однако нас часто искушает желание избежать ответственности, предпочитая, чтобы эти вопросы решались

путём обращения к общественному мнению или к какой-либо форме морального консенсуса или моральной экспертизы. Это искушение можно назвать моральным эскапизмом» (Hertzberg, 2002, p. 252). Философ выделяет два философских заблуждения, подкрепляющих эскапизм. Первое – «вводящая в заблуждение онтология»: идея, что моральные суждения имеют определённый предмет – ценность, добродетель, права, обязанности – и что именно это делает их моральными суждениями (Hertzberg, 2002, p. 253). Если я не знаю, какая столица Перу, я обращаюсь к энциклопедии. Если я не знаю, правильно ли поступаю, кажется, что можно пойти к «моральному эксперту». Однако, как настаивает Гертцберг, моральная неуверенность устроена иначе. Когда я сомневаюсь, правильно ли бросить друга в беде, я сомневаюсь не в том, чего я не знаю, а в том, что я уже понимаю, потому что два моих понимания конфликтуют. Никакой эксперт не может разрешить этот конфликт вместо меня (Hertzberg, 2002, p. 256). Второе заблуждение – «вводящая в заблуждение эпистемология»: представление, что моральный вопрос – это вопрос о фактах, которые можно вычислить (Hertzberg, 2002, p. 254). Философ уверен, что человек, действительно оказавшийся перед моральной дилеммой, никогда не находится в состоянии «чистого листа». Человек уже связан обещаниями, уже испытывает сострадание, уже чувствует вину. Вопрос не в том, что правильно вообще, а в том, как мне быть с этими конкретными людьми здесь и сейчас (Hertzberg, 2002, p. 258).

Однако в постфеноменологической концепции человек не является радикально независимым субъектом, обитающим в мире разрозненных инструментальных объектов; вместо этого субъект и объект совместно формируются в своих опосредованных отношениях (Verbeek, 2011). С этой точки зрения, наши восприятия, практики и даже этика строятся на взаимоотношениях между человеком и технологиями. Так, Дон Иде (Ihde, 1979), американский философ, основатель постфеноменологии, доказал в своих исследованиях, что технологии не являются нейтральными инструментами, а активно опосредуют человеческое восприятие и действие. Именно он впервые вводит понятия «технологического опосредования» и «мультистабильности». Мультистабильность означает, что одна и та же технология может иметь множество различных применений и значений в зависимости от контекста и способа использования. Как пишет Иде, «способы, которыми технология опосредует человеческий опыт и отношения к миру, изменчивы: одна технология может содержать в себе множество применений, влияний и значений, зависящих от того, как она используется отдельным человеком или коллективом, а не только от намерений дизайнера» (Ihde, 1999, p. 46).

А вот нидерландский философ Питер-Пауль Вербек (Verbeek, 2005), продолжая развивать идеи Иде в направлении моральной философии технологий утверждает, что технологии не только опосредуют восприятие, но и активно участвуют в формировании моральных решений. Технологии обладают моральной агентивностью – они могут приглашать к определённым действиям и запрещать другие (Verbeek, 2005, p. 80-85). Он полагает, что «субъект и объект совместно формируются в своих опосредованных отношениях» (Verbeek, 2005, p. 45). Технологии не просто инструменты в руках человека; они со-формируют то, кем человек является и как он действует (Verbeek, 2011, p. 15-18). Так в эмпирических исследованиях технологического опосредования американского философа Роберта Розенбергера (Rosenberger, 2014, p. 375-378) с парковыми скамейками: эти технологии могут быть специально разработаны для того, чтобы предложить пешеходам место, где можно ненадолго присесть, но бездомные могут использовать их в качестве импровизированных кроватей, скейтбордисты – в качестве трамплина, а обычные прохожие – как возвышение для завязывания шнурков.

Идеи Боргманна и вышеперечисленных представителей применяет английский исследователь Эндрю Зельни (Zelny, 2025) для анализа влияния ИИ на практическую мудрость. Он утверждает, что «фронезис – это технологически опосредованная способность, опосредование которой лучше всего понимать с постфеноменологической точки зрения» (Zelny, 2025, p. 3).

Следует отметить, что фронезис, или практическая мудрость, – центральное понятие аристотелевской этики. В «Никомаховой этике» Аристотель (Aristotle, 2000, NE VI.8 1142a) определяет фронезис как способность правильно судить о поступках в конкретной ситуации, принимая во внимание все её обстоятельства. Фронезис достигается посредством опыта и практики, а не через теоретическое обучение. Зельни (Zelny, 2025, p. 8-9) вводит также понятие «активно-интенциональной позиции» – позиции внимательного, умелого и вдумчивого взаимодействия с технологиями в рамках какой-либо практики. Когда человек занимает эту позицию, он критически взаимодействует с технологиями и использует их для более глубокого погружения в деятельность, а не для отстранения от неё. В этой связи уместно использовать дефиницию «моральная дисквалификация», которая была введена американским философом Шеннон Валлор. «Моральная дисквалификация» – процесс, в котором утрачиваются навыки морального суждения из-за того, что они больше не применяются на практике (Vallor, 2015, p. 108). Валлор пишет о том, что без возможностей применять практическую мудрость в реальных условиях мы рискуем атрофировать нашу способность применять практическую мудрость, когда это необходимо (Vallor, 2015). Философ анализирует, как системы ИИ, берущие на себя принятие решений, могут привести к атрофии фронезиса. Она отмечает, что «само собой разумеется, что мало опыта, тем более мудрости, можно развить, если люди работают независимо над разрозненными частями проблемы, оторванными от более широкого контекста социального смысла» (Vallor, 2017, p. 172).

В этом плане интересны исследования Дэвида Золлера, подчёркивающего, что «квалифицированное участие в любой практике открывает уголки реальности» (Zoller, 2017, p. 81) таким образом, каким автоматизация может их закрыть. Эту идею можно усилить, введя понятие «технологического доминирования» – ситуация, когда технология оказывает доминирующее влияние на пользователя, позволяя ему занимать более подчинённую позицию (Sutton, Arnold, Holt, 2018, p. 18-20). Но в то же время в исследованиях Эйсиковиц и Фельдман показано, как ИИ «подрывает нашу способность к практическому суждению» (Eisikovits, Feldman, 2022, p. 182).

Практическая значимость настоящего исследования определяется совокупностью полученных результатов, каждый из которых обладает самостоятельной утилитарной ценностью и может быть непосредственно применен в деятельности по проектированию, оценке и нормативно-правовому регулированию систем искусственного интеллекта. Во-первых, эксплицированный в работе принцип Черчилля в интерпретации Альберта Боргманна представляет собой операционализированный критерий этической оценки ИИ-систем. Данный критерий позволяет выявлять, в какой степени автоматизация той или иной практики сохраняет или, напротив, устраняет человеческое внимание к значимым аспектам деятельности. Рекомендуется использовать этот принцип в процедурах этического аудита алгоритмов, а также при разработке корпоративных стандартов ответственного ИИ, особенно в сферах, где автоматизация сопряжена с риском утраты человеческого суждения (медицина, юриспруденция, военное дело, социальная работа). Во-вторых, выявленные в исследовании формы морального эскапизма применительно к практике использования ИИ позволяют диагностировать ситуации, в которых пользователь систематически делегирует моральную ответственность алгоритму, избегая собственного практического суждения. Рекомендуется использовать данную типологию при проектировании пользовательских интерфейсов и агентных систем, а также в юридической практике для разграничения ответственности между разработчиком, пользователем и самой системой. В-третьих, построенная и обоснованная в статье типология технологического опосредования фронезиса предоставляет практический инструмент для сравнительной оценки различных архитектур ИИ. Рекомендуется применять данную типологию при выборе между альтернативными проектными решениями. Типология может быть адаптирована для использования в отраслевых стандартах по направлению ИИ. В-четвертых, представленный в исследовании критический анализ доминирующей парадигмы проектирования ИИ, фиксирующий систематический приоритет вытесняющих и ослабляющих отношений, имеет прямое практическое значение для разработчиков и инженерных команд. Данный результат позволяет осознанно пересматривать проектные решения, направленные на полную автоматизацию морально значимых выборов, и заменять их архитектурами, сохраняющими агентность пользователя. Рекомендуется включать соответствующие диагностические критерии в чек-листы этической экспертизы проектов на этапах концептуального проектирования и прототипирования.

Таким образом, сформулированные в исследовании принципы «реальной этики ИИ» представляют собой итоговый конструктивный результат, интегрирующий все предшествующие аналитические шаги. Данные принципы задают нормативные ориентиры для проектирования таких систем искусственного интеллекта, которые не замещают, а расширяют интеллектуальные способности человека. Рекомендуется использовать эти принципы: а) при разработке корпоративных кодексов этики в области ИИ; б) в образовательных программах по философии технологий, этике искусственного интеллекта и responsible AI; в) в качестве основы для нормативных документов, регулирующих разработку и внедрение ИИ на отраслевом и национальном уровнях. Таким образом, практическая значимость исследования заключается в том, что его результаты – от диагностических критериев (принцип Черчилля, типология морального эскапизма, классификация опосредования фронезиса) до конструктивных нормативных принципов – образуют готовый инструмент для трансформации существующей парадигмы проектирования ИИ в направлении расширения, а не подавления человеческой моральной агентности.

Обсуждение и результаты

Обсуждение этики искусственного интеллекта начну, опираясь на логику американского постфеноменолога А. Боргманна к искусственному интеллекту. Алгоритмы, интерфейсы и системы поддержки принятия решений – это наши новые «здания». Интерфейс чат-бота, который предлагает готовый ответ на любой вопрос, подобен телевизору в центре гостиной. Система рекомендаций, которая безостановочно подсовывает следующий ролик, подобна автомагистрали, по которой нас везут, пока мы спим. Поэтому в эпоху ИИ принцип У. Черчилля понимается так: мы проектируем наши алгоритмы, а затем наши алгоритмы проектируют нас (Borgmann, 2006, p. 5). Как показано в теоретической части, А. Боргманн различает «вещь» и «устройство». Современные ИИ-системы проектируются преимущественно как устройства. Они предлагают готовые ответы, готовые изображения, готовые решения, а также приглашают к потреблению, а не к практике. Именно поэтому ИИ-системы провоцируют переход от господствующей реальности к одноразовой (Borgmann, 1992, p. 296-297).

Синтез идей философов Гертцберга (Hertzberg, 2002) и Боргманна (Borgmann, 2006) позволяет увидеть, что ИИ является совершенной машиной морального эскапизма. Пользователь, спрашивающий генеративный искусственный интеллект «что мне делать?», находится в том состоянии «чистого листа», которое Гертцберг считает иллюзорным (Hertzberg, 2002, p. 258). Подлинная моральная неуверенность возникает из конфликта уже принятых обязательств, а не из незнания фактов. Перекалывая решение на ИИ, пользователь совершает эскапизм: уклоняется от нравственного выбора.

Нейронные сети, генерируя правдоподобные тексты о добре и долге, симулируют моральное знание, углубляя онтологическое заблуждение – иллюзию, что существует «моральный предмет», о котором можно узнать у эксперта или у большой языковой модели (Hertzberg, 2002, p. 253). Более того, ИИ усиливает эпистемологическое заблуждение: ИИ-система предлагает вычислять моральные решения, как если бы это были задачи оптимизации. Утилитаристский калькулятор, который критикует Гертцберг, сегодня воплощён в алгоритмах, обещающих беспристрастное и оптимальное решение (Hertzberg, 2002, p. 254).

Опираясь на типологию Зельни (Zelny, 2025, p. 6-8), можно выделить четыре типа технологического опосредования фронезиса. Первый тип – вытеснение, когда технология полностью исключает человека из практики,

требующей суждения. Уже сегодня в профессиональной сфере компании заменяют копирайтеров, дизайнеров и сотрудников службы поддержки клиентов полностью автоматизированными системами ИИ (Cooban, 2023). В повседневности приложения генерируют трогательные пожелания для близких, так что пользователь не пишет письмо – его пишет программа. Как отмечает Зельни, «в этих отношениях происходит замещение, если продукт можно получить без практики, то в торговле ничего по-настоящему не теряется» (Zelny, 2025, p. 6), но теряется сама практика. Второй тип – ослабление, когда технология не исключает человека полностью, но сводит его роль к нажатию кнопки «одобрить» или «отклонить». В профессиональной сфере платформы автоматизируют процесс найма, и HR-менеджер остаётся в контуре, но его роль – утвердить решение, которое уже приняло программное приложение. Исследования показывают развитие «предвзятости автоматизации»: человек перестаёт доверять себе и начинает доверять алгоритму, даже когда тот ошибается (Goddard, Roudsari, Wyatt, 2014, p. 370-372). Как пишет Валлор, «без возможностей применять практическую мудрость в реальных условиях мы рискуем атрофировать эту способность» (Vallor, 2015, p. 110). Третий тип – расширение, когда технология не даёт готового ответа, а помогает человеку лучше подумать. В профессиональной сфере разработаны фреймворки интеллектуальной помощи в принятии решений, которые отказываются от явных предложений и вместо этого фокусируются на том, чтобы пользователи запрашивали объяснения, которые могли бы помочь им в собственных размышлениях (Schemmer, Kühl, Satzger, 2021, p. 4-5). Например, приложения для осознанности, такие как Headspace (приложение для медитации), как показывают исследования, улучшают способность к децентрации и эмоциональной регуляции – ключевые компоненты фронезиса (Gavrilova, Zawadzki, 2023, p. 2240-2242; Emmerik, Berings, Lancee, 2018, p. 190-193). Четвёртый тип – определение, когда технология создаёт совершенно новую ситуацию, для которой у нас нет готовых моральных шаблонов. Примером служит анонимность в Интернете. До появления онлайн-форумов анонимное общение было редкостью. Теперь человек может принимать любую онлайн-личность.

Стремительное проникновение генеративных нейросетей в профессиональные сферы ставит под вопрос устоявшиеся представления о характере труда, ценности человеческой экспертизы и самой структуре практического разума. Зельни пишет, что «вскоре понятие практической мудрости будет включать в себя умение грамотно рассуждать в мире, где вокруг нас действуют нечеловеческие силы» (Zelny, 2025, p. 9). Эта диагностика фиксирует фундаментальный сдвиг: человеческий агент более не является единственным источником продуктивных решений – он вынужден выстраивать свои действия в координации с автономно функционирующими алгоритмами. Наиболее наглядные трансформации происходят сегодня в творческих индустриях, которые еще недавно считались недостижимыми для автоматизации. Характерный пример – компания Wizards of the Coast, издатель культовых игр Magic: The Gathering и Dungeons & Dragons. В 2023 году издательство было уличено в использовании изображений, сгенерированных нейросетью Midjourney, для иллюстрирования официальных книг (Good, 2023). До внедрения генеративных моделей компания регулярно нанимала десятки профессиональных художников, каждый из которых получал оплату за свою уникальную работу. С появлением доступных инструментов синтеза изображений экономическая логика изменилась радикально: художник как штатная единица или постоянный подрядчик оказался не нужен. Вместо него потребовался промпт-инженер – специалист, умеющий формулировать текстовые запросы для получения желаемого визуального результата от нейросети. Данный пример иллюстрирует не просто смену профессий, а более глубокий процесс вытеснения – термин, который в данном контексте приобретает буквальное значение. Художник ранее занимал позицию, где его практическая мудрость (фронезис) проявлялась в способности принимать эстетические, композиционные и стилистические решения в процессе создания изображения. Теперь же эта компетенция либо перераспределяется между промпт-инженером и нейросетью, либо редуцируется до оператора, который не столько творит, сколько подбирает удачные формулировки.

В терминологии, предложенной ранее для анализа технологического опосредования фронезиса, этот случай представляет собой классическое вытесняющее отношение: нейросеть не дополняет и не усиливает способности художника, а замещает их. Художник оказывается исключен из производственной цепочки. При этом сам результат – иллюстрация для книги – остается достигнутым, но вопрос о его морально-эстетическом статусе и о последствиях вытеснения человеческой агентности остается открытым. С точки зрения формирующейся «реальной этики ИИ», подобные практики требуют критической оценки. Если следовать диагностике Зельни (Zelny, 2025), практическая мудрость будущего действительно должна включать умение взаимодействовать с нечеловеческими силами. Пример Wizards of the Coast демонстрирует опасный прецедент – рынок уже начинает выбирать первую стратегию, поскольку она экономически эффективна в краткосрочной перспективе. В связи с вышеперечисленным можно согласиться с высказыванием из блога писательницы Джоанны Мациевска: «Я хочу, чтобы ИИ стирал мою одежду и мыл посуду, чтобы я могла заниматься искусством и писательством, а не чтобы ИИ занимался моим искусством и писательством, чтобы я могла стирать и мыть посуду» (Maciejewska, 2024).

Это наблюдение фиксирует перевёрнутую логику технокапитализма: технологии обещали освободить нас от рутины, но начали заменять нас именно в тех практиках, которые составляют ядро человеческого процветания (Vallor, 2017, p. 170-171).

В сфере рекрутинга платформы заявляют, что их системы ИИ могут использоваться на каждом этапе процесса найма: от отбора кандидатов и оценки их профессиональных навыков до собеседования и предложения работы. Исследования показывают, что чрезмерная зависимость от автоматизации снижает уверенность пользователей в собственных способностях и чувство профессиональной автономии (Goddard, Roudsari,

Wyatt, 2014, p. 373). Однако, как отмечают Саттон, Арнольд и Холт (Sutton, Arnold, Holt, 2018, p. 20), возникает проблема технологического доминирования: пользователь подчиняется технологии в процессе принятия решений, что ведёт к снижению квалификации. В случае с рекрутингом это означает, что HR-специалист, который долгое время полагался на ИИ при отборе кандидатов, утрачивает способность самостоятельно оценивать людей – его собственный фронезис атрофируется. Таким образом, выход нейросетей на рынок труда – это не просто технологическая или экономическая проблема. Это проблема пересборки практической мудрости в условиях, когда нечеловеческие силы действуют не только вокруг человека, но и вместо него. Ответ на этот вызов требует не адаптации к вытеснению, а сознательного проектирования таких форм технологического опосредования, при которых нейросети служат расширению, а не отмене человеческого суждения.

В сфере принятия повседневных решений журналист Максвелл Страчан (Strachan, 2023) провёл личный эксперимент, доверив всю свою жизнь ChatGPT, и обнаружил, что это угнетает и нереалистично. Однако другие пользователи социальных сетей рекламируют ChatGPT как «интеллектуальный инструмент для принятия важных и повседневных решений». Один из пользователей разработал подсказку для получения советов, призванную исправить проблему, заключающуюся в том, что ChatGPT предлагал лишь варианты, оставляя в итоге решение за пользователем. Это стремление полностью передать выбор технологии – классический пример морального эскапизма, усиленного верой в алгоритмическую объективность (Hertzberg, 2002, p. 259).

В противовес этим примерам существуют расширяющие технологии. Например, ученые нидерландского Института мозга, познания и поведения Дондерса предложили «машины-рефлексо́ры» – системы, которые не говорят пользователю, что делать, а показывают противоречия в его собственных рассуждениях, предлагают альтернативные гипотезы и задают контрфактические вопросы (Cornelissen, Eerdt, Schraffenberger et al., 2022, p. 2). Эти системы предназначены для «указания на данные, которые могут противоречить принятому решению, или предложения альтернативной гипотезы, основанной на некоторой информации, специфичной для конкретного случая» (Cornelissen, Eerdt, Schraffenberger et al., 2022, p. 2).

Исследования приложений для осознанности показывают, что их использование улучшает показатели, включая наблюдение, описание, действие с осознанностью, безоценочное отношение и не реактивность, а также психологическое и социальное качество жизни (Emmerik, Berings, Lancee, 2018, p. 192-193; Gavrilova, Zawadzki, 2023, p. 2242-2244). Здесь технология работает как «вещь» в терминах Боргманна (Borgmann, 1992, p. 296) – она требует усилия, но это усилие развивает.

Сегодня современная индустрия ИИ отдаёт приоритет вытесняющим и ослабляющим отношениям перед расширяющими. Это объясняется коммерческой логикой: устройства, которые делают работу за пользователя, легче монетизировать, чем инструменты, требующие усилий. Как отмечают Саттон, Арнольд и Холт, «необходимо развивать человеческий опыт, чтобы позволить человеку привнести уникальные возможности в среду принятия решений, не для конкуренции с ИИ, а для дополнения ИИ в процессе совместного принятия решений человеком и компьютером» (Sutton, Arnold, Holt, 2018, p. 22). Особую тревогу вызывает то, что технологии начинают заменять нас именно в тех практиках, которые составляют ядро человеческого процветания. Как утверждает Дэвид Золлер, «квалифицированное участие в любой практике открывает уголки реальности» (Zoller, 2017, p. 81), и делает это таким образом, каким автоматизация может их закрыть. Когда мы позволяем технологиям выполнять квалифицированную задачу за нас, это не только лишает нас возможностей для развития навыков, но и закрывает доступ к определённом способу существования в мире.

А вот концепция «вписывания» Мадлен Акрич (Akrich, 1992) помогает понять, почему доминирующая парадигма проектирования так устойчива. Дизайнеры и технологические компании «вписывают» в технологии определённые сценарии использования, которые делают одни стабильность более вероятными, чем другие (Akrich, 1992, p. 208-210). В случае с ИИ сценарии, ориентированные на эффективность и автоматизацию, получают приоритет над сценариями, ориентированными на рефлексию и развитие навыков.

Однако, как подчёркивают постфеноменологи, ни одна технология не является полностью детерминированной. Мультистабильность оставляет пространство для альтернативных практик использования технологий (Rosenberger, 2014, p. 380-382). Рассматривая ценность расширяющих отношений и активно-интенционального подхода, необходимо понять, как часто и в какой степени нам следует участвовать в этих практиках. Учитывая ограниченное количество часов в сутках и ограниченное количество энергии и внимания, насколько реалистично ожидать от людей активного участия во всех практиках, составляющих их жизнь. Зельни (Zelny, 2025) признаёт этот вопрос. «Хотя для некоторых практик умеренное ослабление или даже вытеснение может быть уместным, – пишет он, – обесценивание активной, целенаправленной позиции в отношении практик, которые нам не особенно интересны, чревато ущербом для способности всех практик быть активно и осмысленно вовлечёнными» (Zelny, 2025, p. 10). Мелкие повседневные дела могут иметь свою собственную цель в укоренении нашей жизни и опыта, служа фундаментом, «чтобы привязать нашу мысль к миру и сформировать нас таким образом, чтобы это способствовало ясности мысли и цели» (Zelny, 2025, p. 10). Даже обыденные задачи, такие как стирка и мытьё посуды, выполняют свою функцию, дополняя полноту нашего жизненного опыта.

Если мы относимся к письму как к досадной помехе, которую можно переложить на ИИ, мы не просто экономим время. Мы обесцениваем саму практику письма. Скользкий путь эскапизма ведёт к тотальному обесцениванию всех практик, требующих усилия (Hertzberg, 2002, p. 260). Поэтому ответ на вопрос о том, должны ли мы быть внимательными ко всему, звучит так: не ко всему одинаково, но к самому принципу внимательности – да. Мы не обязаны быть виртуозами во всём, но мы обязаны не превращать свою жизнь в череду действий, совершаемых «на автопилоте», и с уважением относиться к практикам других людей, даже если сами в них не участвуем.

Реальная этика, по Боргманну (Borgmann, 2006), не отменяет технологий. Она требует ответственности за то, как мы их используем и проектируем (Borgmann, 2006, p. 30). На основе проведённого анализа можно сформулировать несколько принципов «реальной этики ИИ».

Первый принцип – дизайн расширяющих систем. Необходимо требовать от разработчиков, чтобы системы поддержки принятия решений не давали готовых ответов, а задавали вопросы. Примером служат системы IDA (Интеллектуальный анализ документов), которые предлагают объяснения, контрфактические сценарии и альтернативные интерпретации, но не сообщают своего собственного решения (Schemmer, Kühl, Satzger, 2021, p. 4-5). Пользователь должен принять решение сам. Другой пример – структура обсуждения «человек – ИИ», в которой система начинает с выявления мыслей как пользователя, так и системы, затем согласовывает общий язык, и только после этого начинается обсуждение, направленное на выработку более развитой позиции (Ma, Qiaoyi, Xinru et al., 2025, p. 3-4).

Второй принцип – «экономика отдельного домохозяйства» в цифровом мире. Боргманн (Borgmann, 2006, p. 120-121) предлагал начать с малого: с того, где вы поставили телевизор. В эпоху ИИ это означает осознанное управление своей цифровой средой. Важно, где размещён ChatGPT: на рабочем столе, в закладках браузера, в каждом приложении или в отдельной папке, куда пользователь заходит осознанно. Важно, используется ли он для черновиков, которые потом переписываются, или для готовых ответов, которые отправляются без изменений. Эти маленькие решения уже формируют моральную экологию пользователя (Borgmann, 1992, p. 299).

Третий принцип – отказ от эскапизма через образование. Философы и педагоги должны объяснять, что моральный вопрос – не задача на вычисление. Когда пользователь спрашивает ChatGPT «что делать?», он не становится мудрее, он становится зависимее (Hertzberg, 2002, p. 261). Настоящая этическая работа начинается там, где кончаются готовые ответы. Как пишет Гертцберг (Hertzberg, 2002, p. 252), прислушивание к своей совести противопоставляется прислушиванию к другим. Образование должно развивать способность оставаться наедине с моральной неуверенностью и не искать немедленного разрешения у внешних инстанций.

Четвёртый принцип – признание мультистабильности. Постфеноменология учит, что одна и та же технология может быть и устройством, и вещью – в зависимости от позиции пользователя (Ihde, 1999, p. 46; Rosenberger, 2014, p. 375). Нейронные сети могут быть устройством, если пользователь берёт готовые ответы, и вещью, если он использует его как спарринг-партнёра для размышлений, спорит с ним и переформулирует вопросы. Как пишет Эндрю Зельни, «возможность использования или подчинённость реальности, которую позволяет технология, зависит от того, что я называю активно-интенциональной позицией» (Zelny, 2025, p. 8). Реальная этика требует от нас выбора этой позиции.

Пятый принцип – коллективное действие и институциональные изменения. А. Боргманн (Borgmann, 2006, p. 120) признавал, что масштаб институтов ограничивает возможности индивидуальных действий. Изменение этих институтов редко бывает простым. Реальная этика требует не только личных решений, но и коллективной ответственности за дизайн технологий. Мы можем требовать от компаний-разработчиков прозрачности алгоритмов, права на объяснение решений, возможности отключения «умных» функций. Мы можем поддерживать исследования в области объяснимого и человеко-ориентированного ИИ. Мы можем создавать сообщества практиков, которые делятся опытом ответственного использования технологий (Sutton, Arnold, Holt, 2018, p. 22-23).

В проведённом анализе важно также обратиться к параллели между идеями Альберта Боргманна (Borgmann, 2006) и Аллена Бьюкенена (Buchanan, 2002, p. 128), считающего, что фоновые условия, делающие этическую жизнь возможной или невозможной, были упущены из виду. Философ Бьюкенен (Buchanan, 2002, p. 130) призывает к «социальной моральной эпистемологии» – своего рода институционалистскому анализу, выявляющему препятствия для морального развития, например наличие структур, которые способствуют или закрепляют морально опасные ложные убеждения. Он в наибольшей степени подчёркивает то, как институции могут порождать ложные или неполные знания и как эти эпистемологические ошибки могут привести к моральным проступкам (Buchanan, 2002, p. 132-133). В случае с ИИ ложное представление о том, что алгоритмы объективны и беспристрастны, может служить аналогичной функцией: оно легитимирует передачу моральных решений машинам (Eisikovits, Feldman, 2022, p. 185).

Различие между Боргманном и Бьюкененом очевидно: Боргманн (Borgmann, 1992, p. 291) озабочен материальными и технологическими аспектами человеческой жизни, тогда как Бьюкенен (Buchanan, 2002, p. 126-127) занимается нематериальными институциональными структурами, особенно теми, которые связаны с производством и легитимацией знаний. Однако между этими подходами есть несомненное сходство, и их синтез мог бы быть чрезвычайно плодотворным. Материальная среда в виде алгоритмов и интерфейсов и институциональная среда в виде корпоративных практик, правовых норм и образовательных стандартов совместно формируют условия для морального развития или деградации. Ложное убеждение в том, что алгоритмы нейтральны и беспристрастны, не возникает в вакууме – оно культивируется институциональными практиками технологических компаний, маркетинговыми стратегиями и образовательными программами, которые не учат критическому отношению к ИИ (Goddard, Roudsari, Wyatt, 2014, p. 373). Поэтому реальная этика ИИ должна быть одновременно и материальной, и институциональной: она должна обращать внимание и на дизайн интерфейсов, и на социальные практики, которые этот дизайн окружают.

Заключение

В результате проведённого исследования были получены несколько основных выводов, которые можно сформулировать следующим образом. Принцип Черчилля в интерпретации Боргманна полностью применим

к анализу искусственного интеллекта. Алгоритмы, интерфейсы и системы поддержки принятия решений являются теми самыми «зданиями», которые формируют моральные привычки и способности пользователей Доминирующая сегодня парадигма проектирования ИИ создаёт то, что Боргманн называл «одноразовой реальностью»: она сводит практическую мудрость к потреблению, а не к развитию. Различие между «вещью», требующей практики, внимания и навыка, и «устройством», сводящим практику к потреблению, оказалось ключевым аналитическим инструментом для оценки этического влияния технологий. ИИ-системы проектируются преимущественно как устройства, и именно в этом коренится главная этическая проблема. Искусственный интеллект систематически провоцирует моральный эскапизм. Современные технологии подталкивают пользователя к передаче ответственности за моральный выбор алгоритму, подкрепляя те онтологические и эпистемологические заблуждения. ИИ, генерируя правдоподобные тексты о добре и долге, симулирует моральное знание и тем самым углубляет эти заблуждения, создавая иллюзию, что моральные ответы можно вычислить или получить от эксперта. На основе синтеза идей Боргманна и постфеноменологии была подтверждена типология четырёх типов технологического опосредования фронезиса. Вытеснение означает полное исключение человека из практики, требующей суждения. Ослабление сводит роль человека к утверждению или отклонению уже готового решения. Расширение, напротив, поддерживает рефлексию пользователя, не давая готовых ответов. Наконец, определение создаёт новые контексты, требующие совершенно новой мудрости, которой раньше не существовало. Первые два типа доминируют в текущей практике проектирования ИИ, тогда как третий и четвёртый представляют собой альтернативу, соответствующую принципам реальной этики. Моральная дисквалификация – это не теоретический риск, а реально наблюдаемый процесс. Эмпирические исследования подтверждают, что чрезмерная зависимость от автоматизированных систем снижает профессиональные навыки и уверенность в собственных суждениях. Примеры из сферы искусства, где компании заменяют художников генеративным ИИ из сферы рекрутинга, где алгоритмы отбирают кандидатов, а HR-специалисты лишь утверждают решения, и из повседневной практики принятия решений, где люди советуются с ChatGPT по жизненным вопросам, – все они иллюстрируют один и тот же процесс атрофии фронезиса.

На основе проведённого анализа были сформулированы принципы «реальной этики ИИ». Во-первых, необходим дизайн расширяющих систем – таких, которые задают вопросы, а не дают готовые ответы, которые предлагают объяснения и контрфактические сценарии, но оставляют окончательное решение за человеком. Во-вторых, требуется то, что Боргманн называл «экономикой отдельного домохозяйства» – практика осознанного управления своей цифровой средой, внимание к тому, где и как используются ИИ-инструменты. В-третьих, необходимо образование, направленное на развитие фронезиса и критическое отношение к моральному эскапизму; обучающиеся должны понимать, что обращение к ИИ за моральным советом не делает их мудрее, а, напротив, может сделать зависимее. В-четвёртых, следует признать мультистабильность технологий: одна и та же система может быть и устройством, и вещью – в зависимости от того, какую позицию занимает пользователь. Реальная этика требует сознательного выбора активно-интенциональной позиции. В-пятых, необходимы коллективные действия по изменению институциональных условий, включая требования к прозрачности и объяснимости алгоритмов, поддержку исследований в области человеко-ориентированного ИИ и создание сообществ ответственного использования технологий.

Перспективы дальнейших исследований включают несколько направлений. Необходимы эмпирические исследования долгосрочных эффектов использования генеративного ИИ на способность к моральному суждению, особенно в образовательных контекстах, где обучающиеся всё чаще используют нейронные сети для выполнения заданий. Требуется сравнительный анализ разных культурных контекстов использования ИИ и их влияния на фронезис – вполне возможно, что в разных культурных средах эскапизм проявляется по-разному. Важной задачей является разработка и тестирование прототипов «расширяющих» систем поддержки принятия решений, следующих принципам IDA и «машины-рефлекторы». Наконец, представляет значительный интерес дальнейшее исследование возможности синтеза подходов Боргманна, ориентированного на материальную среду, и Бьюкенена, ориентированного на институциональную эпистемологию, для создания комплексной теории морального опосредования технологиями.

Материалы исследования | Research materials

1. Cooban A. This CEO Replaced 90% of Support Staff with an AI Chatbot // CNN. 12.07.2023. <https://edition.cnn.com/2023/07/12/business/dukaan-ceo-layoffs-ai-chatbot>
2. Good O. S. D&D Publisher Updates Policies After AI Art Discovered in Latest Book // Polygon. 07.08.2023. <https://www.polygon.com/23823516/dnd-dungeons-dragons-wizards-ai-art-controversy-bigby-presents-glory-of-the-giants/>
3. Maciejewska J. [@Author]Mac]. You Know What the Biggest Problem with Pushing All-things-AI is? // X. 2024. [https://x.com/Author\]Mac/status/1773679197631701238](https://x.com/Author]Mac/status/1773679197631701238)
4. Strachan M. I Asked Chat-GPT to Control My Life, and It Immediately Fell Apart // Vice. 17.05.2023. <https://www.vice.com/en/article/what-happens-when-you-ask-ai-to-control-your-life/>

Источники | References

1. Akrich M. The De-Scripton of Technical Objects // Shaping Technology / Building Society: Studies in Sociotechnical Change / ed. by W. Bijker, J. Law. MIT Press, 1992.
2. Aristotle. Nicomachean Ethics / ed. by R. Crisp. Cambridge: Cambridge University Press, 2000.

3. Borgmann A. Real American Ethics: Taking Responsibility for Our Country. Chicago: University of Chicago Press, 2006.
4. Borgmann A. The Moral Significance of the Material Culture // Inquiry. 1992. Vol. 35. No. 3-4. <https://doi.org/10.1080/00201749208602295>
5. Buchanan A. Social Moral Epistemology // Social Philosophy and Policy. 2002. Vol. 19. No. 2. <https://doi.org/10.1017/S0265052502192065>
6. Cornelissen N. A. J., Eerdt R. J. M., van, Schraffenberger H. K., Haselager W. F. G. Reflection Machines: Increasing Meaningful Human Control over Decision Support Systems // Ethics and Information Technology. 2022. Vol. 24. No. 2. <https://doi.org/10.1007/s10676-022-09645-y>
7. Eisikovits N., Feldman D. AI and Phronesis // Moral Philosophy and Politics. 2022. Vol. 9. No. 2. <https://doi.org/10.1515/mopp-2021-0026>
8. Emmerik A. A. P. van, Berings F., Lancee J. Efficacy of a Mindfulness-Based Mobile Application: A Randomized Waiting-List Controlled Trial // Mindfulness. 2018. Vol. 9. No. 1. <https://doi.org/10.1007/s12671-017-0761-7>
9. Gavrilova L., Zawadzki M. J. Examining How Headspace Impacts Mindfulness Mechanisms Over an 8-week App-Based Mindfulness Intervention // Mindfulness. 2023. Vol. 14. <https://doi.org/10.1007/s12671-023-02214-4>
10. Goddard K., Roudsari A., Wyatt J. C. Automation Bias: Empirical Results Assessing Influencing Factors // International Journal of Medical Informatics. 2014. Vol. 83. No. 5. <https://doi.org/10.1016/j.ijmedinf.2014.01.001>
11. Hertzberg L. Moral Escapism and Applied Ethics // Philosophical Papers. 2002. Vol. 31. No. 3. <https://doi.org/10.1080/05568640209485105>
12. Ihde D. Technics and Praxis. Dordrecht: Reidel, 1979.
13. Ihde D. Technology and Prognostic Predicaments // AI & Society. 1999. Vol. 13. No. 1-2.
14. Ma S., Qiaoyi C., Xinru W., Chengbo Z., Zhenhui P., Ming Y., Xiaojuan M. Towards Human-AI Deliberation: Design and Evaluation of LLM-Empowered Deliberative AI for AI-Assisted Decision-Making // CHI 2025: CHI Conference on Human Factors in Computing Systems. 2025. <https://doi.org/10.1145/3706598.3713423>
15. Rosenberger R. Multistability and the Agency of Mundane Artifacts: From Speed Bumps to Subway Benches // Human Studies. 2014. Vol. 37. <https://doi.org/10.1007/s10746-014-9317-1>
16. Rosenberger R., Verbeek P.-P. Postphenomenological Investigations: Essays on Human-Technology Relations. Lexington Books, 2015.
17. Schemmer M., Kühl N., Satzger G. Intelligent Decision Assistance versus Automated Decision-Making // arXiv. 2021. <https://doi.org/10.48550/arXiv.2109.13827>
18. Sutton S. G., Arnold V., Holt M. How Much Automation Is Too Much? Keeping the Human Relevant in Knowledge Work // Journal of Emerging Technologies in Accounting. 2018. Vol. 15. No. 2. <https://doi.org/10.2308/jeta-52311>
19. Vallor S. AI and the Automation of Wisdom // Philosophy and Computing / ed. by T. M. Powers. Springer, 2017. Vol. 128. https://doi.org/10.1007/978-3-319-61043-6_8
20. Vallor S. Moral Deskilling and Upskilling in a New Machine Age: Reflections on the Ambiguous Future of Character // Philosophy & Technology. 2015. Vol. 28. <https://doi.org/10.1007/s13347-014-0156-9>
21. Verbeek P.-P. Moralizing Technology: Understanding and Designing the Morality of Things. Chicago: University of Chicago Press, 2011.
22. Verbeek P.-P. What Things Do: Philosophical Reflections on Technology, Agency, and Design. University Park: Penn State University Press, 2005.
23. Zelny A. Offloading Wisdom: Four Technological Relations that Mediate Phronesis // Philosophy & Technology. 2025. Vol. 38. <https://doi.org/10.1007/s13347-025-00889-2>
24. Zoller D. Skilled Perception, Authenticity, and the Case Against Automation. Oxford: Oxford University Press, 2017. <https://doi.org/10.1093/oso/9780190652951.003.0006>

Информация об авторах | Author information



Дедюлина Марина Анатольевна¹, к. филос. н., доц.

¹ Южный федеральный университет, г. Ростов-на-Дону



Marina Anatolyevna Dedyulina¹, PhD

¹ Southern Federal University, Rostov-on-Don

¹ madedyulina@sfedu.ru

Информация о статье | About this article

Дата поступления рукописи (received): 04.05.2026; опубликовано online (published online): 24.06.2026.

Ключевые слова (keywords): реальная этика; моральный эскапизм; расширение фронеизиса; искусственный интеллект; технологическое опосредование; real ethics; moral escapism; expansion of phronesis; artificial intelligence; technological mediation.