

смартфон более 5 часов в день, можно говорить о системном риске для их психофизиологического состояния. Необходим поиск методов коррекции данного состояния, направленных на психологическую разгрузку и на восстановление баланса организма в условиях цифровой среды.

Литература.

1. Авдеева, Е. А. Влияние цифровой электронной среды на когнитивные функции школьников и студентов / Е. А. Авдеева, О. А. Корнилова // КВТиП. – 2022. – № S3. – URL : <https://cyberleninka.ru/article/n/vliyanie-tsifrovoy-elektronnoy-sredy-na-kognitivnye-funktsii-shkolnikov-i-studentov> (дата обращения: 22.03.2026).

2. Влияние цифровой среды на умственную работоспособность и мышление учащихся / Е. С. Богомоллова, К. А. Лангуев, Н. В. Котова, Е. В. Лангуева // Наука и школа.- 2022. - №1. -URL : <https://cyberleninka.ru/article/n/vliyanie-tsifrovoy-sredy-na-umstvennuyu-rabotosposobnost-i-myshlenie-uchaschihsya> (дата обращения: 22.03.2026).

3. The brain in your pocket: evidence that smartphones are used to supplant thinking / N. Barr, G. Pennycook, J. A. Stolz [et al.] // Comput. HumanBehav. – 2015. – Vol. 48. – P. 473–480.

УДК 354

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ: ОТВЕТСТВЕННОСТЬ, АВТОНОМИЯ И БУДУЩЕЕ ТРУДА

Щебелёва Д.И.

Научный руководитель – Климентьева И.А.

УО «Витебская ордена «Знак Почета» государственная академия
ветеринарной медицины», г. Витебск, Республика Беларусь

*В статье анализируются три главные этические проблемы искусственного интеллекта: кто отвечает за действия автономных систем, как сохранить свободу человека и как справедливо решить проблему замены людей на рабочих местах. Доказывается, что для их решения нужен новый подход, объединяющий философию, право и экономику, а также новые правила для регулирования ИИ. **Ключевые слова:** искусственный интеллект (ИИ), доктрина строгой ответственности, распределенная моральная агентность, автономия ИИ.*

ARTIFICIAL INTELLIGENCE: RESPONSIBILITY, AUTONOMY, AND THE FUTURE OF WORK

Shchebeleva D.I.

Scientific supervisor – Klimentyeva I.A.

Vitebsk State Academy of Veterinary Medicine, Vitebsk, Republic of Belarus

*This article analyzes three key ethical issues in artificial intelligence: who is responsible for the actions of autonomous systems, how to preserve human freedom, and how to fairly resolve the problem of replacing humans in the workplace. It argues that addressing these issues requires a new approach that integrates philosophy, law, and economics, as well as new rules for regulating AI. **Keywords:** artificial intelligence (AI), strict liability doctrine, distributed moral agency, AI autonomy.*

Введение. Актуальность данного исследования заключается в феномене искусственного интеллекта (ИИ), который развивается быстрее, чем создаются нормы для его использования. Возникают вопросы: кто отвечает за его действия, как защитить людей, что угрожает правам человека и стабильности общества. Цель данного исследования - анализ этических проблем ИИ: ответственность, автономия и влияние на трудовые отношения. Для реализации поставленной цели были определены задачи: сравнить две модели ответственности за действия ИИ - строгую и распределённую; определить требования к человеческому контролю над ИИ.

Материалы и методы исследований. Концептуальный анализ для переосмысления ключевых понятий «ответственность», «автономия» в контексте понятия ИИ; системный подход для изучения взаимодействий в системе «разработчик – ИИ - общество»; сравнительный анализ этических теорий для оценки моральных дилемм ИИ; нормативный подход для разработки конкретных этических принципов.

Результаты исследований. Главная этическая проблема ИИ - кризис ответственности. Обычные правовые системы требуют найти конкретного виновника - человека, который понимал бы последствия своих действий. Автономные же системы, такие как беспилотный автомобиль, ломают эту логику. Возникает «разрыв ответственности» (responsibility gap): нельзя обвинить разработчика (он не мог все предусмотреть), пользователя (он не управлял процессом) или ИИ (у него нет сознания и умысла). Об этой проблеме писал философ Андреас Матиас в статье «Responsibility Gap and Artificial Intelligence».

Чтобы решить эти вопросы, нужны новые подходы. Один из них - доктрина строгой ответственности, где за ущерб всегда отвечает компания-оператор, которая получает прибыль от технологии. Подобный подход рассматривается в европейских инициативах по регулированию ИИ, например, в «Акте об искусственном интеллекте» ЕС. Другой путь - распределённая моральная агентность, где ответственность делится между всеми участниками: инженерами, регуляторами и пользователями. Концепцию развивала в своих исследованиях профессор Вирджиния Дигнум. Таким образом, вместо поиска виновного мы создаём систему управления рисками, которая предотвращает проблемы и справедливо компенсирует ущерб.

Другая проблема ИИ - конфликт между искусственным и человеческим интеллектом. Самостоятельность ИИ может вступать в конфликт со свободой выбора человека. Это проявляется на разных уровнях. На индивидуальном уровне: алгоритмы соцсетей и рекомендательных систем (как в YouTube или TikTok) создают «пузырь фильтров», показывая только то, что удерживает наше внимание. Это ограничивает нашу когнитивную автономию - способность самостоятельно формировать взгляды. На социальном уровне: алгоритмы принимают или помогают принимать решения в суде, при выдаче кредитов или приёме на работу, они могут невольно воспроизводить и усиливать историческую несправедливость и предрассудки.

Ключевые принципы использования ИИ: объяснимость (XAI) - логика решений ИИ должна быть понятна людям для проверки и оспаривания; контроль - в важнейших вопросах окончательное решение должно оставаться за человеком; трансформация труда - переход от экономической эффективности к этике социальной справедливости так как автоматизация ИИ угрожает не только ручным, но и умственным профессиям. Труд-это не просто заработок, но и способ самореализации и обретения социального статуса. Массовое замещение людей машинами грозит ростом безработицы, неравенства и потерей смысла жизни, что противоречит принципам справедливости. Предлагаются различные модели решения данной проблемы, которые сейчас активно обсуждаются. Это и универсальный базовый доход (UBI) - гарантированная выплата каждому гражданину для обеспечения базовой экономической безопасности (эксперименты проводились в Финляндии, Кении и ряде других стран). Другой вариант – это смещение акцента на «человеческие» сферы, т.е. переориентация экономики на области, где машины пока слабо используются (творчество, уход, образование, управление сложными межличностными отношениями). Если экономическая ценность человеческого труда падает, то акцент следует сделать, на такой деятельности как искусство, наука, волонтерство и общение.

Заключение. Таким образом, развитие ИИ создает глубокие философско-этические проблемы ответственности, автономии и справедливости. Классическая модель поиска одного виновника для случаев причинения вреда автономной системой не работает. Ей на смену должна прийти модель распределенной ответственности и управления рисками. Конфликт между автономией ИИ и человека можно смягчить через обеспечение прозрачности алгоритмов и сохранение окончательного человеческого контроля в критически важных вопросах. Проблема трансформации труда требует не только экономического, но и гуманистического подхода. Технологический прогресс должен сопровождаться перераспределением благ (например, через дискуссию о безусловном базовом доходе) и созданием условий для новой реальности, где ценность человека не сводится только к его профессиональной функции. Преодоление этих проблем зависит не от самих алгоритмов, а от

морального выбора общества, который необходимо сделать уже сегодня через открытый междисциплинарный диалог.

Литература.

1. Матиас, А. Разрыв ответственности: приписывание ответственности за действия самообучающихся автоматов / А. Матиас // Этика и информационные технологии. – 2004. – Т. 6, № 3. – С. 175–183 – [Текст - электронный]. – URL : <https://doi.org/10.1007/s10676-004-3422-1>. – (дата обращения: 21.12.2025).

2. ОЭСР. Принципы ОЭСР в области искусственного интеллекта / Организация экономического сотрудничества и развития. – Париж : ОЭСР, 2019 [Текст - электронный]. – URL : <https://www.oecd.org/digital/ai-principles.htm>. – (дата обращения : 23.10.2024).

3. Планируемое нормативное регулирование искусственного интеллекта: Регламент Европейского парламента и Совета, устанавливающий гармонизированные правила в области искусственного интеллекта (Закон об искусственном интеллекте). // Официальный журнал Европейского союза. – 2024 [Текст - электронный]. – URL : <https://digital-strategy.ec.europa.eu/en/policies/european-ai-act>. – (дата обращения : 21.12.2025).

УДК 141.7

ФУТУРОЛОГИЯ ИОАХИМА ФЛОРСКОГО

Шепилевич А.А.

Научный руководитель – Девярых С.Ю.

УО «Витебская ордена «Знак Почета» государственная академия ветеринарной медицины», г. Витебск, Республика Беларусь

*Иоахим Флорский (ок. 1132–1202) — итальянский теолог и мистик. Его идеи оказали влияние на последующую европейскую мысль, находя отклик как в религиозных движениях, так и в социально-утопических теориях, что позволяет некоторым исследователям считать его одним из предшественников современных философий истории. **Ключевые слова:** Иоахим Флорский, философия истории, средневековая теология, тринитарная концепция, эра Святого Духа, утопия, хилиазм.*

JOACHIM OF FLORUS' FUTUROLOGY

Shepilevich A.A.

Supervisor – Devyatykh S.Yu.

Vitebsk State Academy of Veterinary Medicine, Vitebsk, Republic of Belarus