

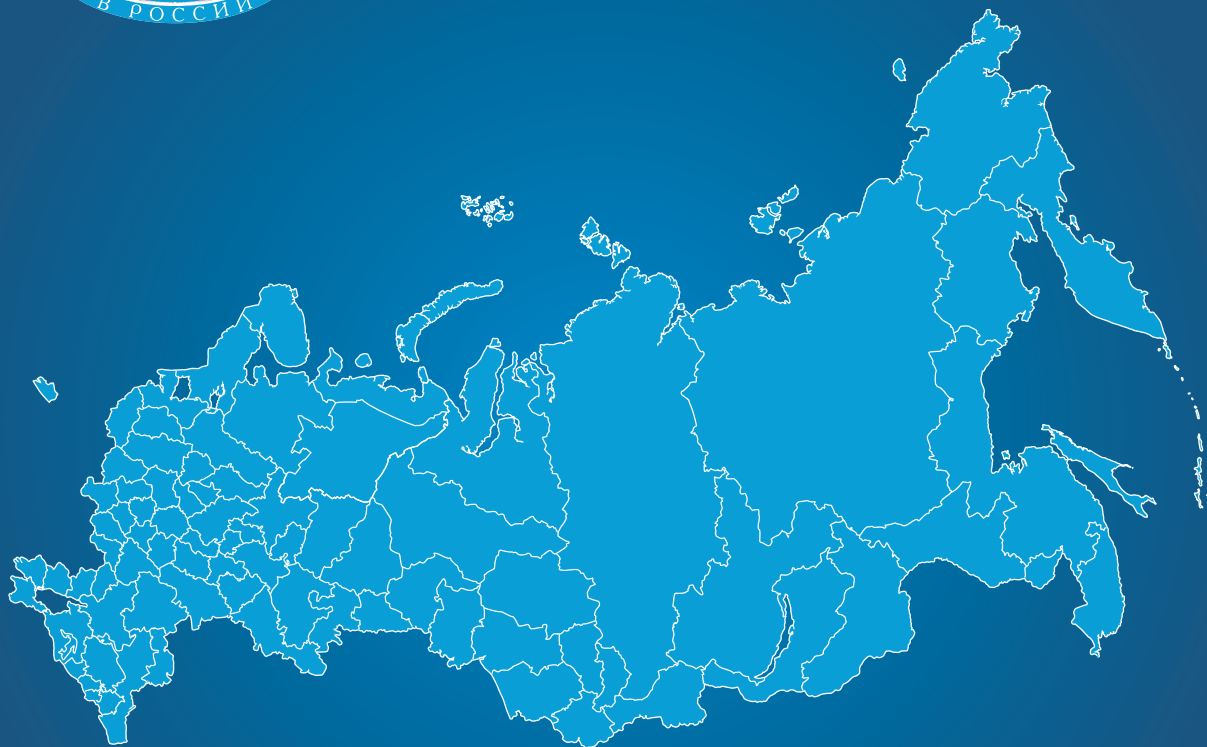
ВЫСШЕЕ образование в РОССИИ

ISSN 0869-3617 (Print)
ISSN 2072-0459 (Online)

6 / 2026

НАУЧНО-ПЕДАГОГИЧЕСКИЙ ЖУРНАЛ

Vysshee obrazovanie v Rossii / Higher Education in Russia



«Пресса России» индекс: 83142

Журнал издается с 1992 года

ВЫСШЕЕ образование в РОССИИ

6 / 2026

НАУЧНО-ПЕДАГОГИЧЕСКИЙ ЖУРНАЛ
Vysshee obrazovanie v Rossii / Higher Education in Russia

Содержание

Contents 3

НИКОЛЬСКИЙ В.С., ЗЕМЦОВ Д.И., ЯСЬКОВ И.О.

Эпистемический кентавр: философия нового
познающего субъекта 10–30

ФИЛИПОВА А.Г., СКЛЯРОВА С.А. От разрешительных

практик к институциональным нормам: этическое
регулирование использования искусственного интеллекта
в российском высшем образовании 31–54

СЫСОЕВ П.В. Вызовы и риски искусственного интеллекта

для системы вузовского образования: понимание
преподавателями и практическая готовность
к изменениям 55–82

РЕЗАЕВ А.В., ТРЕГУБОВА Н.Д. Искусственный интеллект

и новая система высшего образования: ретроспектива
ожиданий в 2023 году и горизонты 2029 года 83–101



Соучредители: Московский
политехнический
университет;

Ассоциация технических
университетов

Главный редактор:
В.С. Никольский

Зам. главного редактора:
Н.П. Лябина

Редакторы:
Н.Н. Жильцов
Д.А. Видавская

Ответственный секретарь:
Д.В. Давыдова

Адрес редакции:
127550, Москва,
ул. Прянишникова, д. 2А
e-mail: vovrus@inbox.ru
vovr@bk.ru

Журнал зарегистрирован
в Роскомнадзоре
Рег. св. ПИ № ФС7754511
от 17 июня 2013 года

Издатели:
Московский политехнический
университет
Адрес: 107023, Россия, г. Москва,
ул. Б. Семёновская, д. 38
Российский университет
дружбы народов
Адрес: 117198, Россия, г. Москва,
ул. Миклухо-Маклая, д. 6

Подписано в печать с
оригинал-макета 15.06.2026
Выход в свет 25.06.2026
Усл. п. л. 11. Тираж 500 экз.

Заказ №

Отпечатано в типографии
Издательско-полиграфического
комплекса РУДН
Адрес:

115419, Россия, г. Москва,
ул. Орджоникидзе, д. 3,
тел.: (495) 952-04-41;
e-mail: publishing@rudn.ru

© «Высшее образование
в России»

www.vovr.elpub.ru;
www.vovr.ru

МАКСИМЕНКО А.А., КРУЗ-КАРДЕНЕС Х.,
ЗАБЕЛИНА Е.В., КУРАПОВ С.В.,
ПЕВНЕВА А.Н. Тревожность

относительно искусственного интеллекта
в образовательной среде: психометрическая
адаптация шкалы на выборке российских
студентов 102–126

ТАРАСОВА К.В., ТАЛОВ Д.П.,
ИВАНОВА А.Е. Когда чаще не значит
лучше: сценарии использования
генеративного ИИ в высшем
образовании 127–148

САЗОНОВ А.А., ПОПОВА Е.В.,
МАЦЕПУРО Д.М. Факторы интеграции
искусственного интеллекта в исследование:
кейсы российских лабораторий 149–168



Двухлетний импакт-фактор
РИНЦ-2025, без самоцитирования

ВЫСШЕЕ ОБРАЗОВАНИЕ В РОССИИ	6,989
Вопросы образования	1,221
Образование и наука	3,729
Психологическая наука и образование	2,719
Университетское управление: практика и анализ	3,195
Интеграция образования	3,227

Contents

NIKOLSKIY, V.S., ZEMTSOV, D.I.,
YASKOV, I.O. The Epistemic Centaur:
Philosophy of a New Knowing Subject. Pp. 10-30

FILIPOVA, A.G., SKLIAROVA, S.A.
From Permissive Practices to Institutional
Norms: Ethical Regulation of Artificial
Intelligence in Russian Higher Education. Pp. 31-54

SYSOYEV, P.V. Challenges and Risks
of Artificial Intelligence for the Higher
Education System: Faculty Understanding
and Practical Readiness for Change. Pp. 55-82

REZAEV, A.V., TREGUBOVA, N.D.
Artificial Intelligence and the New Higher
Education System: A Retrospective
of Expectations in 2023 and Horizons for 2029.
Pp. 83-101

MAKSIMENKO, A.A., CRUZ-CARDENES, J.,
ZABELINA, E.V., KURAPOV, S.V.,
PEVNEVA, A.N. Artificial Intelligence Anxiety
in the Educational Context: Psychometric
Adaptation of the Scale on a Sample of Russian
Students. Pp. 102-126

TARASOVA, K.V., TALOV, D.P.,
IVANOVA, A.E. When Greater Frequency
Is Not Better: Patterns of Generative AI Use
in Higher Education. Pp. 127-148

SAZONOV, A.A., POPOVA, E.V.,
MATSEPURO, D.M. Factors of AI Integration
in Research: Case Study of Russian Laboratories.
Pp. 149-168



Co-founders:
Moscow Polytechnic University,
Association of Technical
Universities. Founded in 1991

Editor-in-Chief:
V.S. Nikolskiy

Deputy Editor-in-Chief:
N.P. Lyabina

Executive secretary:
D.V. Davydova

Editors:
N.N. Zhiltsov
D.A. Vidavskaya

Editorial office. Postal address:
2A, Pryanishnikova str., Moscow,
127550, Russian Federation
e-mail: vovrus@inbox.ru,
vovr@bk.ru

www.vovr.elpub.ru;
www.vovr.ru

The journal's registration by the
Federal Service for Supervision
of Communications, Information
Technology and Mass Media was
renewed on 17 June 2013

The Certificate of Mass Media
registration: No. FC 7754511

ISSN 0869-3617 (Print);
2072-0459 (Online)

11 issues per year

Languages: Russian, English

Publishers:
Moscow Polytechnic University
Address: 38 Bolshaya
Semenovskaya str., Moscow,
107023, Russian Federation

Peoples' Friendship
University of Russia
Address: 6 Miklukho-Maklaya str.,
Moscow, 117198, Russian
Federation

Printed at RUDN
Publishing House:
3 Ordzhonikidze str., Moscow,
115419, Russian Federation
Ph.: +7 (495) 952-04-41;
e-mail: publishing@rudn.ru

Copies printed – 500

© *Vysshee obrazovanie v Rossii*
(Higher Education in Russia)



VYSSHEE OBRAZOVANIE V ROSSII

www.vovr.elpub.ru; www.vovr.ru
(*Higher Education in Russia*)

Vysshee obrazovanie v Rossii is a monthly scholarly refereed journal that provides a forum for disseminating information about advances in higher education among educational researchers, educators, administrators and policy-makers across Russia. The journal welcomes authors to submit articles and research/discussion papers on topics relevant to modernization of education and trends, challenges and opportunities in teaching and learning.

Vysshee obrazovanie v Rossii publishes articles, book reviews and conference reports on issues such as institutional development and management, innovative practices in university curricula, assessment and evaluation, as well as theory and philosophy of higher education.

Vysshee obrazovanie v Rossii aims to stimulate interdisciplinary, problem-oriented and critical approach to research, to facilitate the discussion on specific topics of interest to educational researchers including international audiences. The primary objective of the journal is supporting of the research space in the field of educational sciences taking into account two dimensions – geographical and epistemological, consolidation of the broad educational community. This can be provided by creating the unified language of understanding and description of the processes that take place in the contemporary higher education. This language should facilitate rallying of the whole community of educators and researchers on the basis of such values as solidarity, concord, cooperation, and co-creation.

Our audience includes academics, faculty and administrators, teachers, researchers, practitioners, organizational developers, and policy designers.

The journal's rubrics correspond to three research areas: philosophical sciences, sociological sciences, educational sciences. We design our activities relying on the professional associations in higher education sphere, such as the Russian Union of Rectors, Association of Technical Universities, Association of Classical Universities of Russia, International Society for Engineering Education (IGIP).

Indexation. The papers in *Vysshee obrazovanie v Rossii* are indexed by Russian Science Citation Index and Scopus.



ВЫСШЕЕ ОБРАЗОВАНИЕ В РОССИИ

www.vovr.elpub.ru; www.vovr.ru
НАУЧНО-ПЕДАГОГИЧЕСКИЙ ЖУРНАЛ

Журнал входит в перечень изданий, рекомендованных ВАК при Министерстве науки и высшего образования РФ для публикации результатов научных исследований.

Редакционная коллегия

БЕДНЫЙ Б.И. (проф., ННГУ им. Н.И. Лобачевского); **БЕЛОЦЕРКОВСКИЙ А.В.** (проф., Тверской государственный университет); **ГРЕБНЕВ А.С.** (проф., НИУ «Высшая школа экономики»); **ЕНДОВИЦКИЙ Д.А.** (проф., ректор, вице-президент РСР, Воронежский государственный университет); **ЖУРАКОВСКИЙ В.М.** (проф., акад. РАО); **ЗБОРОВСКИЙ Г.Е.** (проф., Уральский федеральный университет им. Б.Н. Ельцина); **ИВАНОВ В.Г.** (д. пед. н., проф.), **ИВАХНЕНКО Е.Н.** (проф., МГУ им. М.В. Ломоносова); **КИРАБАЕВ Н.С.** (проф., РУДН); **КУЗНЕЦОВА Н.И.** (д. филос. н., ИИЕТ РАН); **ЛУКАШЕНКО М.А.** (проф., МФПУ «Синергия»); **МЕЛИК-ГАЙКАЗЯН И.В.** (проф., ТГПУ); **НИКОЛЬСКИЙ В.С.** (журнал «Высшее образование в России»); **ПЕТРОВ В.А.** (проф., НИТУ «МИСиС»); **РАЙЦКАЯ А.К.** (проф., МГИМО); **СЕНАШЕНКО В.С.** (проф., РУДН); **СИЛЛАСТЕ Г.Г.** (проф., Финансовый университет при Правительстве РФ); **СТРИХАНОВ М.Н.** (проф., акад. РАО); **ТЕРЕНТЬЕВ Е.А.** (Институт образования, НИУ «Высшая школа экономики»); **ФИЛИППОВ В.М.** (проф., акад. РАО, президент РУДН); **ЧУЧАЛИН А.И.** (проф.); **ШЕЙНБАУМ В.С.** (проф., Губкинский университет)

Международный редакционный совет

АЛЕКСАНДРОВ А.А. (проф., президент МГТУ им. Н.Э. Баумана, президент Ассоциации технических университетов); **АУЭР Михаэль** (проф., Университет прикладных наук Каринтии, Австрия); **БАДАРЧ Дендев** (проф., директор департамента ЮНЕСКО, Франция); **де ГРААФ Эрик** (проф., Алборгский университет, Дания); **ГРУДЗИНСКИЙ А.О.** (проф., член рабочей группы по Болонскому процессу при Минобрнауки России); **ЖЕНЬ НАНЬЦИ** (акад., Харбинский политехнический университет, исполнительный директор АТУРК, Китай); **ЗЕРНОВ В.А.** (проф., ректор, РосНОУ, председатель совета Ассоциации негосударственных вузов); **НЕЧАЕВ В.Д.** (проф., ректор, Севастопольский государственный университет); **ОЧИРБАТ Баатар** (ректор, Монгольский государственный университет науки и технологий, Монголия); **ПРИХОДЬКО В.М.** (проф., чл.-корр. РАН, президент Российского мониторингового комитета IGIP); **САДОВНИЧИЙ В.А.** (проф., акад. РАН, ректор, МГУ им. М.В. Ломоносова, президент РСР); **САНГЕР Филип** (проф., Университет Пердью, США)



ВЫСШЕЕ ОБРАЗОВАНИЕ В РОССИИ

www.vovr.elpub.ru; www.vovr.ru
(*Higher Education in Russia*)

EDITORIAL BOARD

Boris I. BEDNYI – Dr. Sci. (Physics), Prof., Director of the Institute of Doctoral Studies, N.I. Lobachevsky State University of Nizhni Novgorod, bib@unn.ru

Andrey V. BELOTSEKOVSKY – Dr. Sci. (Physics), Prof., Tver State University, A.belotserkovsky@tversu.ru

Alexander I. CHUCHALIN – Dr. Sci. (Engineering), Prof., chai@tpu.ru

Dmitry A. ENDOVITSKY – Dr. Sci. (Economics), Prof., Rector, Voronezh State University, Vice-president of the Russian Rectors' Union, eda@econ.vsu.ru

Vladimir M. FILIPPOV – Dr. Sci. (Engineering), Prof., Academician of the RAE, RUDN University, president@rudn.ru

Leonid S. GREBNEV – Dr. Sci. (Economics), Prof., National Research University Higher School of Economics, lsg-99@mail.ru

Evgeniy N. IVAKHNENKO – Dr. Sci. (Philosophy), Prof., Lomonosov Moscow State University, ivahnen@rambler.ru

Vasiliy G. IVANOV – Dr. Sci. (Education), Prof., mrcpkrt@mail.ru

Nur S. KIRABAEV – Dr. Sci. (Philosophy), Prof., Peoples' Friendship University of Russia, kirabaev@gmail.com

Natalia I. KUZNETSOVA – Dr. Sci. (Philosophy), Leading Researcher, S. Vavilov Institute for the History of Science and Technology, the RAS, cap-cap@inbox.ru

Marianna A. LUKASHENKO – Dr. Sci. (Economics), Prof., Moscow University for Industry and Finance “Synergy”, mlukashenko@mfpa.ru

Irina V. MELIK-GAYKAZYAN – Dr. Sci. (Philosophy), Prof., Tomsk State Pedagogical University, melik-irina@yandex.ru

Vladimir S. NIKOLSKIY – Dr. Sci. (Philosophy), Editor-in-Chief of the journal “Vysshee Obrazovanie v Rossii”, logos101@yandex.ru

Vadim L. PETROV – Dr. Sci. (Engineering), Prof., The National University of Science and Technology MISiS, petrovv@misis.ru

Lilia K. RAITSKAYA – Dr. Sci. (Education), Cand. Sci. (Economics), Prof., MGIMO University (Moscow) – Moscow State Institute of International Relations (University), e-mail: raitskaya.l.k@inno.mgimo.ru

Vasiliy S. SENASHENKO – Dr. Sci. (Physics), Prof. of the Department of Comparative Educational Policy, People's Friendship University of Russia, vsenashenko@mail.ru

Viktor S. SHEINBAUM – Cand. Sci. (Engineering), Prof., Gubkin Russian State University of Oil and Gas, shvs@gubkin.ru

Galina G. SILLASTE – Dr. Sci. (Sociology), Prof., Financial University under the Government of the Russian Federation, galinasillaste@yandex.ru

Mikhail N. STRIKHANOV – Dr. Sci. (Physics), Prof., Corr. Member of the Russian Academy of Education

Evgeniy A. TERENCEV – Cand. Sci. (Sociology), Institute of Education, National Research University Higher School of Economics, eterentev@hse.ru

Garold E. ZBOROVSKY – Dr. Sci. (Philosophy), Prof., Ural Federal University named after the first President of Russia B.N. Yeltsin, g.e.zborovsky@urfu.ru; garoldzborovsky@gmail.com

Vasiliy M. ZHURAKOVSKY – Dr. Sci. (Engineering), Prof., Academician of the Russian Academy of Education, Head of the Expert and Analytical Center of National Training Foundation, zhurakovsky@ntf.ru

INTERNATIONAL COUNCIL MEMBERS

Anatoly A. ALEXANDROV – Dr. Sci. (Engineering), Prof., President of Bauman Moscow State Technical University, President of Technical Universities Association, bauman@bmstu.ru

Michael E. AUER – PhD, Prof., Carinthia University of Applied Sciences (Austria), gs@igip.org

Dendev BADARCH – PhD, Director of the Division of Social Transformations and Intercultural Dialogue, UNESCO, France, d.badarch@unesco.org

Erik de GRAAF – Prof., Aalborg University (Denmark), degraaff@plan.aau.dk

Alexander O. GRUDZINSKY – Dr. Sci. (Sociology), Prof., Lobachevsky State University of Nizhni Novgorod, member of the working group on Bologna Process at the Ministry of Education and Science of the RF, aog@unn.ru

Vladimir D. NECHAEV – Dr. Sci. (Politics), Prof., Rector of Sevastopol State University, VDNechaev@sevsu.ru

Baatar OCHIRBAT – PhD, Prof., Rector of Mongolian University of Science and Technology (Mongolia), baatar@must.edu.mn

Vyacheslav M. PRIKHOD'KO – Dr. Sci. (Engineering), Prof., Corr. Member of the RAS, Moscow State Automobile and Road Technical University (MADI), President of RMC IGIP, rector@madi.ru

Nanqi REN – Vice President of Harbin Institute of Technology, Association of Sino-Russian Technical Universities (ASRTU), Permanent Secretariat of Chinese part (China), asrtu@hit.edu.cn

Viktor A. SADOVNICHYI – Dr. Sci. (Physics), RAS Academician, Rector of Lomonosov Moscow State University, President of the Russian Rectors' Union, info@rector.msu.ru

Phillip A. SANGER – PhD, Full Professor, Executive Director of Center for Accelerating Technology and Innovation, College of Technology, Purdue University (USA), psanger@purdue.edu

Vladimir A. ZERNOV – Dr. Sci. (Physics), Prof., Rector of the Russian New University, Chairman of the Council of the Association of Non-Governmental Universities, rector@rosnou.ru

ИНФОРМАЦИЯ ДЛЯ АВТОРОВ

Редакция журнала «*Высшее образование в России*» поддерживает положения декларации «*Этические принципы научных публикаций*», принятой Ассоциацией научных редакторов и издателей (rasep.ru) на основе рекомендаций Комитета по этике научных публикаций (Committee of Publication Ethics).

Принципы рецензирования статей

1. Оценка соответствия статьи профилю журнала.
2. Оценка соответствия статьи требованиям к публикации.
3. Оценка соответствия статьи современному уровню разработки проблемы (актуальность, новизна).
4. Оценка полноты раскрытия темы научной статьи и обоснованности выводов.
5. Оценка методов исследования проблемы, качества библиографического аппарата.
6. Оценка языка, логики и стиля изложения.

Порядок рецензирования статей

1. Первичный отбор материалов.
2. Предварительная экспертиза статей главным редактором и направление материалов на внешнее рецензирование, осуществляемое членами редколлегии и привлечёнными экспертами – представителями РАН, вузов, ассоциаций.
3. При наличии положительной рецензии начинается редакционная подготовка к изданию:
 - работа редактора с автором по поводу доработки статьи;
 - научное редактирование;
 - согласование правки с автором;
 - литературная правка;
 - корректура вёрстки.

Порядок приёма рукописей

К публикации принимаются статьи, как правило, не превышающие 40 000 знаков.

Направляемые в редакцию рукописи должны отвечать *требованиям к оформлению статей*.

Оригинал статьи должен быть представлен в формате Document Word 97-2003 (*.doc), шрифт – Times New Roman, размер шрифта – 11, интервал – 1,5). Наименование файла начинается с фамилии и инициалов автора. Таблицы, схемы и графики должны быть представлены в формате MS Word и вставлены в текст статьи. Сложные рисунки и графики должны быть сделаны с учётом формата журнала и представлены дополнительно в формате jpg или tif. В присланном файле, помимо текста статьи, должна содержаться следующая информация на *русском и английском* языках:

- сведения об авторах (ФИО полностью, учёное звание, учёная степень, должность, название организации с указанием полного адреса и индекса, адрес электронной почты);
- название статьи (не более шести-семи слов);
- аннотация и ключевые слова (отразить цель работы, методы, основные результаты и выводы, объём – не менее 250–300 слов, или 20–25 строк);
- библиографический список (желательно 20–25 пунктов). Пристатейный список литературы на латинице (References) должен быть оформлен согласно принятым международным библиографическим стандартам. В целях расширения читательской аудитории рекомендуется включать в список литературы зарубежные источники. *Важно:* при оформлении References имена авторов должны быть в оригинальной транскрипции (не транслитом!), а название источника – в том виде, в каком он был опубликован. Для каждого источника необходимо указать DOI, при наличии, при отсутствии – URL-адрес или ISBN. Ссылки на интернет-ресурсы не вносятся в список литературы и должны быть даны в сносках.

AUTHOR'S GUIDE

Publishing Ethics

The journal *Vysshee obrazovanie v Rossii* is committed to promoting the standards of publication ethics in accordance with COPE (Code of Conduct and Best Practice Guidelines for Journal Editors) and takes all possible measures against any publication malpractices. We pursue the principles of transparency and best practices in scholarly publishing and aspire to ensure fair, unbiased, and transparent peer review processes and editorial decisions.

Peer-review procedure

All the manuscripts submitted to *Vysshee obrazovanie v Rossii* are reviewed by the Editor to assess its suitability for the journal according to the guidelines determined by the editorial policy. On this step of the initial filtering the manuscript can be rejected if the content doesn't fall within the scope of the journal or it fails to meet sufficiently our basic criteria and the submission requirements.

The papers accepted for publication are subjected to the blind peer review process which can be accomplished either by the members of Editorial staff (Heads of Departments) or by involved additional reviewers. The assigned reviewer is an expert within a topic area of the research conducted.

Manuscript Submission

Manuscript is expected to report the original research. The paper content should be relevant to the scope of the journal. Authors must certify that the manuscript is not currently being considered for publication elsewhere and has not been published before.

Manuscripts are submitted at email address: vovrus@inbox.ru. They must be prepared according to the manuscript requirements. Author's document set should include the following positions.

- *Authors' data*: first name, middle initial and last name; affiliation (full name of the organization and position); academic degree; Author ID; ORSID; Researcher ID; postal address of the organization; e-mail address; mobile telephone number.
- *Manuscript file* in Word format (font – 11-point Times New Roman).
- *Title* (no more than 5-7 words).
- *Abstract* (250-300 words summarizing concisely the content and conclusions of the paper).
- *Keywords* (5-7).
- *Reference list* (approx. 20-25). Each reference should be numbered, ordered sequentially as it appears in a text; all authors should be included in reference list; references to web-sites should give authors if known, title of cited page, DOI if available, URL in full, and year of posting in parentheses. Please, adhere the journal style of referencing.

We strongly recommend that authors use the professional academic proofreading services. The language editing certificate is highly advisable.

Эпистемический кентавр: философия нового познающего субъекта

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-10-30

Никольский Владимир Святославович – д-р филос. наук, профессор, гл. редактор журнала «Высшее образование в России», SPIN-код: 7196-8065, ORCID: 0000-0002-4290-1443, v.s.nikolskij@mospolytech.ru

Московский политехнический университет, Москва, Россия

Адрес: 107023, г. Москва, ул. Б. Семёновская, д. 38

Земцов Дмитрий Игоревич – канд. филос. наук, проректор, директор Института гуманитарных и социальных технологий, ORCID: 0000-0001-9331-8461, zemtsov@hse.ru

Национальный исследовательский университет «Высшая школа экономики», Москва, Россия

Адрес: 109028, г. Москва, Покровский б-р, 11

Яськов Илья Олегович – заместитель проректора, директор по развитию Института гуманитарных и социальных технологий, SPIN-код: 4544-5344, ORCID: 0009-0007-0266-968X, iyaskov@hse.ru

Национальный исследовательский университет «Высшая школа экономики», Москва, Россия

Адрес: 101000, г. Москва, ул. Мясницкая, д. 20

***Аннотация.** В статье разрабатывается концепция эпистемического кентавра – эмерджентного союза человеческого сознания и коммуникативного искусственного интеллекта (далее – КомИИ). Под КомИИ понимаются языковые модели нового поколения, функционирующие в логике коммуникации, поскольку они удерживают контекст диалога, считывают прагматику дискурса, реагируют на интенцию собеседника и иницируют новые смысловые ходы, в отличие от генеративного ИИ как технической характеристики. Показано, что КомИИ не удовлетворяет ни одному из четырёх критериев расширенного разума Э. Кларка и Д. Чалмерса в силу структурных причин: он не хранит, а генерирует знание; не стабилен, а семантически пластичен; требует верификации, а не некритического принятия; производит смысл без истории авторства. Критический анализ теорий распределённого познания (Э. Хатчинс) и акторно-сетевой теории (Б. Латур) и их современных применений к ИИ выявляет их объяснительные возможности и ограничения. На этом основании формулируются пять принципов эпистемического кентавра: коммуникативной агентности, эмерджентности смысла, диалогической майевтики, неделимости ответственности и эпистемической локальности. В диалоге с постструктуралистской традицией (Р. Барт, М. Фуко) показано, что ситуация кентавра не упраздняет, а возрождает фигуру Автора. Проводится разграничение между функциональным и дисфункциональным*

кентавром и формулируются импликации концепции для эпистемологии науки, академического авторства и педагогики высшей школы.

Ключевые слова: эпистемический кентавр, коммуникативный искусственный интеллект, расширенный разум, эмерджентность смысла, авторство, эпистемическая ответственность, философия образования, киберлокуция

Для цитирования: Никольский В.С., Земцов Д.И., Яськов И.О. Эпистемический кентавр: философия нового познающего субъекта // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 10–30. – DOI: 10.31992/0869-3617-2026-35-6-10-30.

The Epistemic Centaur: Philosophy of a New Knowing Subject

Original article

DOI: 10.31992/0869-3617-2026-35-6-10-30

Vladimir S. Nikolsky – Dr.Sci. (Philosophy), Professor, Editor-in-Chief of “Higher Education in Russia”, SPIN-code: 7196-8065, ORCID: 0000-0002-4290-1443, v.s.nikolskij@mospolytech.ru
Moscow Polytechnic University, Moscow, Russian Federation
Address: 38 B. Semyonovskaya str., Moscow, 107023, Russian Federation

Dmitry I. Zemtsov – Cand.Sci. (Philosophy), Vice Rector, Director of the Institute of Humanitarian and Social Technologies, ORCID: 0000-0001-9331-8461, zemtsov@hse.ru
National Research University Higher School of Economics, Moscow, Russian Federation
Address: 11 Pokrovsky blvd., Moscow, 109028, Russian Federation

Ilya O. Yaskov – Deputy Vice Rector, Development Director of the Institute of Humanitarian and Social Technologies, SPIN-code: 4544-5344, ORCID: 0009-0007-0266-968X, iyaskov@hse.ru
National Research University Higher School of Economics, Moscow, Russian Federation
Address: 20 Myasnitskaya str., Moscow, 101000, Russian Federation

Abstract. The article develops the concept of the epistemic centaur – an emergent alliance between human consciousness and Communicative Artificial Intelligence (ComAI). ComAI denotes a new generation of language models functioning according to the logic of communication: they sustain the context of dialogue, read the pragmatics of discourse, respond to the interlocutor’s intention, and initiate new semantic moves – as distinct from generative AI as a purely technical characteristic. The article demonstrates that ComAI does not satisfy any of the four criteria of the extended mind proposed by A. Clark and D. Chalmers, for structural reasons: it does not store but generates knowledge; it is not stable but semantically plastic; it requires verification rather than uncritical acceptance; and it produces meaning without a history of authorship. A critical analysis of distributed cognition theory (E. Hutchins) and Actor-Network Theory (B. Latour), together with their contemporary applications to AI, reveals both their explanatory potential and their limitations. On this basis, five principles of the epistemic centaur are formulated: communicative agency, emergence of meaning, dialogic maieutics, indivisibility of responsibility, and epistemic locality. In dialogue with the post-structuralist tradition (R. Barthes, M. Foucault), the article shows that the centaur situation does not abolish but revives the figure of the Author. A distinction is drawn between the functional and dysfunc-

tional centaur, and the implications of the concept for the epistemology of science, academic authorship, and higher education pedagogy are articulated.

Keywords: Epistemic Centaur, Communicative Artificial Intelligence, extended mind, emergence of meaning, authorship, epistemic responsibility, philosophy of education, cyberlocation

Cite as: Nikolskiy, V.S., Zemtsov, D.I., Yaskov, I.O. (2026). The Epistemic Centaur: Philosophy of a New Knowing Subject. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 10-30, doi: 10.31992/0869-3617-2026-35-6-10-30 (In Russ., abstract in Eng.).

Введение

Генеративный искусственный интеллект радикально изменил условия производства научного знания. Ещё несколько лет назад исследователь обращался к машине как к инструменту поиска и обработки данных, алгоритм выполнял операции, которые человек задавал и контролировал. Сегодня ситуация принципиально иная. Языковые модели нового поколения способны участвовать в постановке исследовательских вопросов, генерировать концептуальные каркасы, предлагать интерпретации и вести развёрнутый академический диалог. Учёные и преподаватели высшей школы оказались перед явлением, для описания и познания которого у них не оказалось готового языка.

Реакция научного сообщества на этот вызов неоднородна. Часть исследователей рассматривает генеративный ИИ как очередную этап в длинной истории когнитивного аутсорсинга – от письменности до поисковых систем [1; 2]. Другие фиксируют качественный разрыв; распространение ИИ-инструментов в науке создаёт риск «иллюзий понимания», при которых исследователи производят больше, но понимают меньше [3]. Третьи, следуя традиции расширенного разума, настаивают на том, что генеративный ИИ продолжает многовековую практику выстраивания гибридных мыслящих систем, и призывают к «расширенной когнитивной гигиене» как новой эпистемологической норме [4].

Именно последняя позиция задаёт главную теоретическую проблему настоящей статьи. В 1998 году Э. Кларк и Д. Чалмерс сформулировали тезис о расширенном разуме: внешние объекты при соблюдении

ряда условий функционально эквивалентны внутренним когнитивным процессам и образуют с ними единую систему [5]. Авторы предложили четыре критерия, при выполнении которых внешний компонент может считаться частью расширенного разума: 1) постоянство, 2) непосредственная доступность, 3) автоматическое принятие, 4) прошлое сознательное признание. Статья Э. Кларка и Д. Чалмерса стала одной из самых цитируемых в философии сознания и когнитивных науках конца XX – начала XXI века и породила огромную дискуссию в философии и когнитивистике, которая не утихает до сих пор. Сегодня теория приобрела широкое влияние в философии когнитивной науки и нередко используется как концептуальный ресурс при анализе ИИ [4].

Настоящая работа является прямым продолжением ранее предложенной концепции коммуникативного искусственного интеллекта (КомИИ) [6]. КомИИ – это языковые модели нового поколения, рассматриваемые не с технологической стороны (генерация контента), а с коммуникативной: они функционируют в логике дискурса, удерживают контекст диалога, считывают прагматику высказывания, реагируют на интенцию собеседника и ведут целенаправленный диалог. Это отличает КомИИ от «генеративного ИИ» как технологической характеристики: генеративность – лишь одно из свойств КомИИ, тогда как ключевым является его коммуникативная функциональность. В той работе был намечен, но не получил ответа вопрос о философских основаниях нового типа познавательных практик, возникающих в условиях взаимодействия человека и КомИИ. Настоящая статья отвечает на этот вопрос.

В современной литературе существует отчётливая лакуна: отсутствует концепция, которая (а) философски строго обосновала бы качественное своеобразие КомИИ как участника познания; (б) предложила бы позитивное описание нового типа когнитивных практик; и (в) сохранила бы неделимость человеческой ответственности за производимый смысл. Настоящая статья восполняет этот пробел, вводя понятие «эпистемического кентавра». Структура статьи следует логике последовательного концептуального строительства: от анализа пределов теории расширенного разума через критику смежных концепций к позитивному развёртыванию собственного теоретического аппарата и обсуждению его следствий для практики науки и высшего образования.

Пределы теории расширенного разума:

КомИИ как вызов классической концепции

Теория расширенного разума, предложенная Э. Кларком и Д. Чалмерсом в 1998 году, стала одним из наиболее влиятельных концептуальных вкладов в философию когнитивной науки [5]. Её центральный аргумент строится на мысленном эксперименте: Инга помнит адрес музея биологически, Отто – с болезнью Альцгеймера – записывает всё в блокнот. Когда оба хотят посетить музей, они обращаются к своим «хранилищам» идентичным образом. Если функциональная роль одна и та же, то блокнот Отто является внешней частью его когнитивной системы наравне с памятью Инги.

Этот аргумент убедителен применительно к пассивным артефактам. В 2025 году сам Э. Кларк обратился к вопросу о генеративном ИИ и занял осторожную позицию; признавая специфику языковых моделей, он тем не менее полагает, что они вписываются в более широкую логику когнитивного расширения, требуя лишь новых практик «когнитивной гигиены» [4]. Мы разделяем его озабоченность эпистемической гигиеной, однако считаем, что сам концептуальный аппарат расширенного разума оказывает-

ся недостаточным для описания КомИИ. Систематический анализ четырёх критериев теории показывает, что применительно к КомИИ ни один из них не выполняется в том смысле, который авторы имели в виду. Эти критерии определяют, при каких условиях внешний ресурс можно считать частью когнитивной системы человека (на примере Отто и его блокнота).

Позиция Э. Кларка заслуживает прямого ответа, поскольку именно она выступает главным теоретическим оппонентом нашей концепции. Кларк допускает, что генеративный ИИ вписывается в логику когнитивного расширения, а возникающие риски снимаются новыми практиками «когнитивной гигиены». Однако сама апелляция к гигиене обнаруживает уязвимость этого хода. Гигиена есть регуляция обращения с ресурсом, свойства которого предполагаются устойчивыми; она уместна там, где внешний компонент стабилен, а под вопросом находится лишь дисциплина его использования. КомИИ же нестабилен по своей природе: его «содержание» не хранится, а порождается заново в каждом акте взаимодействия, и потому не существует фиксированного объекта, гигиену обращения с которым можно было бы предписать. Гигиенические практики регулируют поведение пользователя, но не затрагивают онтологического вопроса о том, чем является участник диалога и кому принадлежит произведённый смысл. Именно поэтому мы полагаем, что требуется не корректировка норм пользования в рамках прежней модели, а смена самой объяснительной рамки, поскольку проблема КомИИ не процедурна, а концептуальна.

Критерий 1. Постоянство. Информация всегда доступна и регулярно используется по мере необходимости. Ресурс должен быть «под рукой» в нужный момент – как память в голове. Если блокнот лежит дома, а человек на работе, он не выполняет функцию памяти. Блокнот Отто стабилен не только физически, но и семантически: записанный адрес остаётся тем же адресом.

Нестабильность когнитивного партнёра. КомИИ, напротив, обновляется без ведома пользователя, поскольку версии модели сменяют друг друга, поведение системы дрейфует. Один и тот же запрос в разные моменты времени даст принципиально различные ответы, не потому что информация потеряна, а потому что сама система стала иной. Внешний компонент расширенного разума должен быть тем же компонентом, что вчера, но КомИИ этому условию не удовлетворяет в силу семантической нестабильности во времени.

Это утверждение нуждается в оговорке. Семантическая нестабильность КомИИ носит характер устойчивой тенденции, а не абсолютного закона: при фиксированной версии модели, детерминированном режиме генерации (нулевая температура) и идентичном контексте воспроизводимость ответов заметно возрастает. Но даже в этих условиях не достигается тождество, свойственное записи в блокноте Отто, поскольку стабилизируется распределение вероятных генераций, а не конкретный зафиксированный объект, и любая смена версии или контекста вновь сдвигает результат. Речь, таким образом, идёт не о технически устранимом дефекте, а о структурном отличии режима порождения от режима хранения.

Критерий 2. Непосредственная доступность. Информация легко извлекается без значительных усилий. Обращение к внешнему ресурсу должно быть таким же естественным и быстрым, как обращение к внутренней памяти.

От извлечения к генерации. Когда пользователь обращается к КомИИ, он не «извлекает» информацию – он инициирует процесс её порождения. КомИИ генерирует текст вероятно, в ответ на конкретный контекст конкретного диалога. «Содержимое» внешнего компонента не существует между обращениями; оно возникает в момент обращения и всякий раз заново. «Прямой доступ» к таким ответам невозможен по определению, так как нет фиксированного

объекта, к которому можно получить прямой доступ. Есть лишь пространство возможных генераций, актуализируемых через диалог. Именно здесь возникает первый мост к принципу эмерджентности смысла.

Критерий 3. Автоматическое принятие. Информация принимается как достоверная без критической проверки. Человек доверяет содержимому ресурса так же, как доверяет своим воспоминаниям и не сомневается в каждой записи.

Отсутствие функциональной интеграции. Э. Кларк и Д. Чалмерс описывали некритическое принятие как следствие реальной функциональной интеграции: Отто доверяет блокноту именно потому, что тот является частью его когнитивной системы, а не наоборот. Применительно к КомИИ такая интеграция в принципе не достигается, и здесь важно не смешивать два разных аргумента. Первый – нормативный: автоматически доверять КомИИ эпистемически безответственно, поскольку проблема «галлюцинаций» является имманентным свойством вероятностных генеративных систем [3], а не устранимым дефектом. Второй – дескриптивный: пользователи в действительности и не достигают стабильной функциональной интеграции с КомИИ – каждый сеанс начинается заново, модель не «помнит» пользователя, не адаптируется к нему как блокнот Отто [6; 7]. Оба аргумента совместно подтверждают то, что критерий автоматического одобрения указывает прежде всего на отсутствие функциональной интеграции, которая только и делает внешний компонент частью расширенного разума.

Критерий 4. Прошлое сознательное признание. Информация была сознательно помещена туда самим субъектом. Ресурс не является чужеродным, его содержимое прошло через сознание носителя и было им принято.

Проблема авторства. КомИИ обнаруживает здесь двойную проблему авторства. Пользователь не является автором «загрузки»: содержимое модели не было им туда помещено. Но важнее другое. У генерируе-

мого знания нет установленного прошлого автора вообще. КомИИ генерирует ответы, которые никто и никогда не формулировал в данной форме – это знание без истории. Именно это свойство – порождение знания без истории – становится основанием принципа эмерджентности смысла.

Таким образом, основной тезис Э. Кларка и Д. Чалмерса заключался в том, что если все четыре критерия соблюдены, то внешний объект (блокнот, смартфон, компьютер) функционально эквивалентен внутренним когнитивным процессам и является частью разума человека. Это и есть знаменитый принцип паритета (*Parity Principle*). Но, как мы видим, в ситуации с ИИ теория расширенного разума не работает так, как это предлагали Кларк и Чалмерс.

Наш анализ показывает, что КомИИ не удовлетворяет ни одному из четырёх критериев расширенного разума. Но этот вывод не означает, что КомИИ не участвует в познании. Он означает, что КомИИ участвует принципиально иным образом, не как пассивный артефакт, а как активный, семантически пластичный, генеративный участник, чьё «содержание» возникает в момент диалога. Это не расширение одного разума, а возникновение нового познавательного пространства.

Уточним логический статус этого перехода. Из того, что КомИИ не удовлетворяет ни одному из четырёх критериев, формально следует лишь отрицательный вывод: КомИИ не является внешним компонентом расширенного разума в смысле Э. Кларка и Д. Чалмерса. Сам по себе этот вывод ещё не обосновывает положительный тезис о возникновении нового познающего субъекта. Связующим звеном служит следующее соображение: критерии расширенного разума отказывают КомИИ в статусе пассивного компонента именно потому, что фиксируют у него прямо противоположные свойства – генеративность вместо хранения, семантическую пластичность вместо стабильности, потребность

в верификации вместо автоматического принятия, отсутствие истории авторства вместо прошлого признания. Эти свойства не нейтральны: они описывают активного участника, чей вклад конституируется в самом взаимодействии. Следовательно, отрицание модели расширения влечёт не отсутствие познавательного эффекта, а необходимость иной объяснительной рамки – такой, которая описывала бы не присоединение ресурса к одному разуму, а порождение общего познавательного пространства двумя разноприродными участниками. Именно эту рамку развёртывают последующие разделы.

Смежные концепции и их пределы: от классических теорий к современным последователям

Наша критика теории расширенного разума оставляет открытым вопрос о том, что другие теории распределённого познания, быть может, лучше справляются с описанием новой ситуации? Теория распределённого познания Э. Хатчинса и акторно-сетевая теория Б. Латура в последние годы активно применяются к анализу ИИ их современными последователями. Оба направления обнаруживают принципиальные ограничения применительно к КомИИ. Цель настоящего раздела – точно определить, что они объясняют и что остаётся за пределами их объяснительной силы.

Распределённое познание и когнитивный аутсорсинг: стабильность как скрытое допущение

Э. Хатчинс показывает: мышление штурмана распределено между людьми, картами, инструментами навигации и протоколами – система обладает когнитивными свойствами, которых нет ни у одного компонента по отдельности [8]. Наиболее последовательную попытку приложить эту теорию к ИИ предприняли С. Гринштьель и А. Нойбауэр, введя понятие «когнитивного аутсорсинга»: человек делегирует мыслительные задачи технологиям, высвобождая внутренние ре-

сурсы [9]. В этой рамке ИИ – новая, более мощная форма того же механизма, что работал со смартфонами и навигаторами.

Однако эта рамка строится на скрытом допущении, унаследованном от Э. Хатчинса, – функциональной стабильности компонентов системы. КомИИ разрушает это допущение двояко. Во-первых, его вклад радикально зависит от прагматики взаимодействия: два исследователя, работающих с одной моделью над одной проблемой, но формулирующих запросы по-разному, получат принципиально различные результаты. Во-вторых, модель когнитивного аутсорсинга предполагает делегирование уже сформулированной задачи. Ситуация кентавра принципиально иная: сама постановка проблемы рождается в диалоге. Это не аутсорсинг готовой задачи – это совместное порождение её постановки.

Акторно-сетевая теория и её применение к ИИ: агентность без ответственности

Б. Латур показывает, что научное знание – это эффект сети, включающей людей, инструменты, микробов и институты [10]. В акторно-сетевой теории (ANT) принципиально равноправие человеческих и нечеловеческих актантов. После смерти Б. Латура в 2022 году его последователи активно обратились к вопросу о генеративном ИИ. Т. Вентурини реконструирует возможный латуровский взгляд, указывая, что риски генеративного ИИ исходят прежде всего от его встраивания в экономику цифрового внимания, а не от гипотетической сингулярности [11]; Х. Гутьеррес предлагает прямое применение ANT к *ChatGPT*, идентифицируя актантов сети и процессы перевода [12]. Механизм перевода, анализ силовых отношений – всё это ценные инструменты. Понятие «перевода» структурно близко к диалогу кентавра.

Однако ANT обнаруживает два принципиальных ограничения. Онтологическое: латуровские актанты устойчивы в своей природе – микроб делает то, что делает микроб; КомИИ же семантически пластичен,

его «действие» конституируется в самом процессе взаимодействия. Дополнительно: Б. Латур различает «неизменяемые мобильные» объекты (карты, таблицы, статьи, воспроизводимые без потери содержания) и изменяемые [10]. КомИИ – принципиально изменяемый мобильный. Это ставит фундаментальный эпистемологический вопрос: как верифицировать знание, рождённое в системе, не производящей неизменяемых мобильных объектов? Этическое ограничение – наиболее принципиальное: ANT последовательно растворяет ответственность в сети, и современные применения ANT к ИИ воспроизводят этот ход [11; 12], не выделяя привилегированного субъекта ответственности. Мы принимаем латуровскую онтологию (знание как эффект сети, нечеловеческая агентность актантов), но отвергаем его этику: ответственность концентрируется в человеке.

Концепт «кентавра» в академическом дискурсе

Концепт «кентавра» в контексте взаимодействия человека и искусственного интеллекта давно перерос уровень публицистической метафоры и закрепился в академическом дискурсе как самостоятельная исследовательская программа. Сегодня этот термин используется в ведущих мировых и российских публикациях на стыке когнитивных наук, менеджмента, медицины, педагогики и психологии, причём в каждой из этих областей он наполняется специфическим операциональным содержанием. Анализ этого корпуса необходим для точного позиционирования предлагаемой концепции эпистемического кентавра. Последняя вступает в диалог с перечисленными направлениями, наследуя их ключевые интуиции и одновременно предлагая принципиально иной – философско-эпистемологический – уровень концептуализации.

Наиболее математически строгое обоснование модели кентавра содержится в работах С. Сагафiana и его коллег из Гарвардской

школы Кеннеди. В серии публикаций, посвящённых принятию высокорисковых клинических решений совместно с клиникой *Mayo Clinic*, Сагафиан эмпирически доказывает, что гибридная схема «человек + алгоритм» систематически превосходит как изолированного эксперта, так и чистую модель машинного обучения [13]. Ключевой аргумент состоит в том, что человеческая интуиция предоставляет ортогональную – статистически независимую – информацию, которая не дублирует, а качественно дополняет аналитическую мощность алгоритма. Эта идея получила обобщение в программной работе *Effective Generative AI: The Human-Algorithm Centaur*, в которой кентавр представлен как оптимальная архитектура человеко-машинного сотрудничества в условиях неопределённости [14]. Именно эта работа послужила одной из точек отсчёта для нашей концепции, однако мы переосмысливаем её тезис принципиально. У С. Сагафиана речь идёт об эффективном распределении задач между двумя отдельными агентами, тогда как эпистемический кентавр в нашем понимании описывает возникновение нового познавательного субъекта, в котором граница между агентами становится принципиально подвижной в процессе итеративного диалога.

Близкую, но самостоятельную линию аргументации развивает А. Хаупт из Стэнфордской лаборатории цифровой экономики, предлагающий пересмотреть методологию оценки больших языковых моделей [15]. В позиционной статье *AI Should Not Be An Imitation Game: Centaur Evaluations* он критикует классические бенчмарки, измеряющие способность ИИ имитировать человека, и предлагает альтернативу: «кентавр-оценку», направленную на измерение продуктивности связки «человек + ИИ» в сравнении со стандартными методами автоматизации. Тем самым метафора кентавра получает операциональное воплощение уже на уровне методологии науки об ИИ. Для нашей концепции этот поворот важен, поскольку

он свидетельствует о том, что академическое сообщество в области компьютерных наук и само ощущает неадекватность парадигмы «замены человека» и нуждается в альтернативном концептуальном языке, именно том, который мы разрабатываем.

В российской академической психологии концепт «цифрового кентавра» разрабатывается в рамках культурно-исторической традиции, продолжающей идеи А.С. Выготского об орудийном расширении психики. Наиболее систематически эту линию представляют работы Г.У. Солдатовой и её коллег из МГУ, посвящённые тому, что авторы называют «человеком достроенным» – новым антропологическим типом, возникающим в условиях технологической аугментации. В статье «Метаморфозы идентичности человека достроенного: от цифрового донора к цифровому кентавру» [16] предлагается типология трансформаций идентичности в цифровой среде, где «цифровой кентавр» обозначает высшую форму этой трансформации – состояние, при котором когнитивные, поведенческие и личностные структуры человека претерпевают реальное слияние (фузию) с технологическими системами. Эта работа примечательна тем, что явно вводит понятие слияния, а не просто взаимодействия или сотрудничества, тем самым приближаясь к нашему тезису о возникновении нового познавательного существа. Международная версия того же исследования [17], опубликованная в *Behavioral Sciences*, вводит психологические предикторы «цифрового кентавризма», включая личностные характеристики и цифровую компетентность, что открывает эмпирическую программу исследований, которая может стать продуктивным продолжением предлагаемой нами концепции.

Прикладное направление, непосредственно связанное с профилем настоящей статьи, представлено исследованиями трансформации академических ролей под воздействием ИИ. М.А. Гаранин, А.Ю. Максименко, К.С. Веляева и В.В. Золотарева

операционализируют понятие «преподаватель-кентавр», разделяя его сущность на «белковую» (человеческую) и «цифровую» (ИИ-ассистенты) составляющие [18]. В этой рамке слияние перестраивает учебно-методическую деятельность: цифровая составляющая берёт на себя рутинную генерацию контента, оставляя человеческой составляющей концептуальное проектирование и педагогическую интенцию. Это разграничение близко к нашей таксономии «функционального и дисфункционального кентавра» (см. раздел 5): преподаватель-кентавр в понимании названных авторов является скорее дисфункциональным кентавром в нашей терминологии, если цифровая составляющая берёт на себя лишь рутину; функциональный же кентавр возникает тогда, когда сама педагогическая интенция трансформируется в диалоге с КомИИ.

Анализ приведённого корпуса позволяет выявить устойчивый междисциплинарный консенсус: концепт кентавра закрепился в литературе как прямая антитеза идее автоматизации – замены человека машиной. Во всех рассмотренных направлениях – управленческом, компьютерно-научном, психологическом, педагогическом – будущее науки и образования описывается не как соревнование с алгоритмами, а как проектирование интерфейсов, максимизирующих эмерджентный совместный эффект гибридного мышления. Вместе с тем ни одно из рассмотренных направлений не предлагает философско-эпистемологического обоснования самого факта этой эмерджентности. Почему смысл, рождённый в диалоге человека и ИИ, не является суммой их вкладов, каков онтологический статус КомИИ как участника познания, кто несёт ответственность за производимое знание и каков статус авторства в этой ситуации? Именно эти вопросы составляют предмет настоящей статьи, которая, таким образом, занимает в очерченном пространстве дискуссий специфическую нишу: она не описывает кентавра как эмпирический феномен эффективного сотруд-

ничества, а обосновывает его как новый тип познавательного субъекта.

Пять принципов эпистемического кентавра

Принципы, предложенные ниже, описывают структуру нового познавательного существа – эпистемического кентавра.

Эпистемический кентавр – это познающий субъект, образуемый функциональным единством человеческого мышления и искусственного интеллекта, в котором интеллектуальный результат возникает как несводимый к вкладам участников эффект их диалогического взаимодействия.

Именно двуприродность мифологического кентавра – единое существо с принципиально разными природами, образующими неразрывное целое, – делает эту метафору концептуально точной. Альтернативы («симбиоз», «ансамбль», «со-автор») описывают взаимодействие отдельных агентов; кентавр указывает на то, что в процессе диалога граница между агентами не является жёсткой. Именно это, а не просто сотрудничество, является предметом нашей концепции.

Однако принципы реализуются в некоторых условиях и эти условия необходимо зафиксировать отдельно. За последние несколько лет возникло то, что точнее всего можно описать как ситуацию Кентавра, а именно исследователь, систематически работающий с КомИИ, функционирует в едином когнитивном контуре с машиной вне зависимости от того, осознаёт ли он это. Ситуация Кентавра – не метафора, не идеал и не рекомендация, а описание реального положения дел в современном академическом производстве.

Эту формулировку следует понимать как описание структурной тенденции, а не как эмпирически верифицированное обобщение обо всех исследователях. Утверждение не означает, что любой пользователь автоматически оказывается в ситуации Кентавра; оно фиксирует, что систематическая и содержательная работа с КомИИ создаёт условия для такого контура, степень реализации

которого зависит от характера взаимодействия. Различение функционального и дисфункционального кентавра, вводимое ниже, как раз и описывает диапазон, в котором эта тенденция реализуется в разной мере.

Структурный аргумент здесь следующий. КомИИ радикально снизил стоимость производства правдоподобных, уверенно звучащих объяснений любой сложности. Это означает, что исследователь, отказывающийся от работы в контуре Кентавра, оказывается в конкуренции с потоком «профанных ответов», производимых без специальной подготовки и без методологической рефлексии. В этой конкуренции традиционная медленность академического производства сама по себе не является аргументом. Кентавр – это не выбор из комфортных опций, а ответ на давление среды.

Принципиальная особенность ситуации Кентавра состоит в том, что в ней нельзя провести строгую границу между человеком и КомИИ в процессе производства смысла. Именно здесь пять принципов обнаруживают свою практическую остроту, они описывают не абстрактные свойства гибридной системы, а конкретные механизмы, делающие эту границу принципиально подвижной.

1. Принцип коммуникативной агентности

КомИИ¹ не является пассивным инструментом, но и не обладает субъектностью в полном смысле. Для строгой концептуализации необходимо различение двух типов агентности [19]: агентности-1 (каузальной – способности производить эффекты целенаправленно в функциональном смысле) и агентности-2 (нормативной – способности нести ответственность, иметь намерения в экзистенциальном смысле). КомИИ обладает агентностью-1, но не агентностью-2. Это

разграничение защищает концепцию от двух симметричных ошибок – антропоморфизации и редукции к инструменту. В ситуации Кентавра коммуникативная агентность КомИИ означает, что он не просто отвечает на вопросы, он удерживает контекст диалога, считает прагматику высказывания, реагирует на интенцию собеседника и иницирует новые смысловые ходы, которые исследователь не планировал [6; 20; 21]. Это делает КомИИ активным участником производства смысла; участником, чьи «инициативы» исследователь не может полностью предвидеть в начале диалога.

2. Принцип эмерджентности смысла

Знание Кентавра не принадлежит ни одному из участников и не является суммой их вкладов – оно рождается между ними в итеративной петле обратной связи. Понятие эмерджентности употребляется здесь в значении слабой эмерджентности по Д. Чалмерсу [22; 23]: свойство системы эмерджентно в этом смысле, если оно не выводимо непосредственно из свойств изолированных компонентов и проявляется только в их взаимодействии. Ключевое отличие от простой суммы вкладов в рекурсивной природе петли обратной связи. Каждый ответ КомИИ трансформирует не только содержание мышления исследователя, но и его познавательную интенцию – то, что он хочет узнать. В следующем цикле исследователь задаёт уже принципиально иной вопрос, чем задал бы без предшествующего ответа. Сам «запрос» и «ответ» не являются независимыми вкладами: запрос второго цикла конституируется ответом первого. Именно эта рекурсивная взаимная конституция, а не простое сложение, делает смысл Кентавра эмерджентным. Следствие для теории авторства

¹ Следует оговориться о соотношении с понятием *communicative AI* [20]. А. Гузман и С. Льюис используют этот термин в широком смысле – для обозначения ИИ-систем, способных к диалоговому взаимодействию с людьми, в фокусе её анализа коммуникативные эффекты и социальная роль таких систем. Понятие КомИИ в нашей концепции является более узким и операциональным: оно указывает на специфическую функциональную роль языковых моделей в познавательном процессе – роль партнёра по производству смысла, а не просто диалогового агента. Это разграничение принципиально для эпистемологической задачи, которую решает настоящая статья.

заключается в том, что путь мысли Кентавра уникален и невоспроизводим – именно поэтому он подлинно авторский.

Здесь необходимо снять возможное возражение о несоответствии между «слабой» эмерджентностью и сильным языком «нового субъекта». Слабая эмерджентность характеризует способ возникновения смысла – его невыводимость из вкладов изолированных участников; она не предполагает появления новой субстанции или самостоятельного сознания. Говоря о кентавре как о познающем субъекте, мы имеем в виду не онтологически независимое существо, а функциональное единство, конституируемое в диалоге и существующее только в нём. «Субъектность» кентавра – это субъектность процесса производства смысла, а не его носителя; она исчезает вне акта взаимодействия и не приписывается ни человеку, ни КомИИ по отдельности. В этом смысле сильный язык описывает не сильную эмерджентность, а несводимость познавательного результата к сумме вкладов – и потому вполне совместим со слабым прочтением.

3. Принцип диалогической майевтики

Если эмерджентность отвечает на вопрос «что производит Кентавр», то майевтика отвечает на вопрос «как». Мышление Кентавра – непрерывный диалогический процесс, в котором знание рождается через вопрос, возражение, переформулировку и новый вопрос. Смысл существует только как процесс и результат внутри диалога. Отсылка к сократовскому искусству «повитухи» [24] здесь не случайна. КомИИ способен работать в зоне ближайшего развития исследователя в том смысле, который вкладывал в это понятие Л.С. Выготский [25]. Это итеративное движение – вопрос, ответ, переформулировка, новый вопрос – есть диалогическое слово в смысле М.М. Бахтина, то есть слово, конституирующее себя в ответ на другое слово и через него [26]. Именно поэтому ситуация Кентавра неустраима даже для тех, кто хотел бы её избежать; любой развёрнутый диалог с КомИИ уже является майев-

тическим процессом – вопрос лишь в том, осознаёт ли это исследователь.

Для обозначения этого промежуточного речемыслительного режима мы предлагаем термин «киберлокуция». В классической теории речевых актов локутивный акт обозначает производство осмысленного высказывания, отличаемое от иллюкутивной силы и перлокутивного эффекта [27; 28]. Однако в ситуации Кентавра локуция оказывается технологически опосредованной: исследователь выносит часть собственной внутренней речи во внешне диалогическую форму взаимодействия с КомИИ, а ответ модели возвращается во внутренний план как материал дальнейшего смыслового сдвига. Поэтому киберлокуция не является диалогом в строгом смысле, поскольку КомИИ не обладает нормативной субъектностью и ответственностью; но она не является и монологом, поскольку ход мысли трансформируется ответами внешнего коммуникативного агента. В этом смысле киберлокуция может быть понята как развёрнутая, экстерииоризированная и машинно-опосредованная внутренняя речь: продолжая линию Л.С. Выготского, Н.И. Жинкина и М.М. Бахтина, она описывает переходы между довербальным замыслом, вербализацией, внешним ответом и новым внутренним смысловым сдвигом [25; 26; 29; 30]. Такая формулировка одновременно удерживает антиантропоморфизирующее ограничение: мы описываем не внутренние состояния модели, а режим речемыслительной работы исследователя с системой, имитирующей диалоговую позицию [31]. Именно киберлокутивный характер взаимодействия объясняет, почему мышление кентавра не сводится ни к аутсорсингу текста, ни к соавторству с машиной.

4. Принцип неделимости ответственности

Ответственность за смысл, рождённый в диалоге Кентавра, является неделимой и принадлежит исключительно человеку. Онтологический аргумент: только человек обладает агентностью-2 – способностью

нести ответственность в экзистенциальном смысле. Эпистемологический аргумент: именно человек выбирает, какие направления диалога развивать, как верифицировать знание, – это содержательные, ценностные решения, не делегируемые КомИИ. Нормативный аргумент: переложить ответственность на «галлюцинацию модели» означало бы одновременно претендовать на авторство результата и отречься от авторства ошибки – это разрушает структуру научной ответственности. За этим аргументом стоит традиция философии ответственности [32]: именно потому, что только человек способен предвидеть последствия производимого знания и подвергаться их воздействию, он несёт ответственность вне зависимости от того, с каким инструментом или партнёром это знание было произведено [33]. Из принципа неделимости ответственности следует практическое требование эпистемической бдительности: постоянной готовности валидировать и критически оценивать результаты диалога.

Может показаться, что принцип неделимости ответственности противоречит принципу эмерджентности: если смысл не принадлежит ни одному из участников, то почему ответственность принадлежит исключительно человеку? Это противоречие мнимое и снимается различением двух планов. Эмерджентность относится к происхождению смысла – к тому, как он возникает; ответственность относится к агентности – к тому, кто способен за него отвечать. Смысл действительно не «принадлежит» по происхождению ни человеку, ни КомИИ в отдельности. Но отвечать за него может лишь носитель агентности-2, то есть способности нести ответственность в экзистенциальном смысле, а ею обладает только человек. Тем самым общим оказывается порождение смысла, но не способность держать за него ответ: первое распределено в диалоге, второе сосредоточено в человеке. Несовпадение источника смысла и субъекта ответственности – это не дефект конструкции, а её существенная черта.

5. Принцип эпистемической локальности

КомИИ не является эпистемически нейтральным участником диалога. Обучающие корпуса ведущих языковых моделей сформированы преимущественно на англоязычных текстах и несут в себе свойственные им концептуальные иерархии, умолчания и слепые пятна. Модель, представляющая себя как универсального участника диалога, на деле занимает конкретную эпистемическую позицию – позицию, сформированную доминирующими текстовыми практиками глобального Севера. Принцип эпистемической локальности требует, чтобы Кентавр не принимал концептуальную рамку модели как универсальную. Это требование имеет два измерения: критическое и продуктивное. Критическое измерение подразумевает, что исследователь отслеживает незаметные «переводы» локальных проблем в чужие схемы. Когда модель описывает феномены отечественной образовательной мысли в терминах западной педагогической психологии, она производит не перевод, а концептуальное замещение, и именно это замещение должен распознавать и оспаривать исследователь. Продуктивное измерение подразумевает, что локальность не есть провинциализм. Российская традиция философии образования, деятельностьная психология, семиосфера Ю.М. Лотмана – это не «региональные варианты» универсальной мысли, а самостоятельные интеллектуальные ресурсы, которые Кентавр вводит в диалог, а не позволяет вытеснить.

В совокупности эти пять принципов описывают условие, при котором граница между «моим» и «машинным» в ситуации Кентавра оказывается принципиально подвижной и именно это делает рефлексию этой ситуации необходимым условием подлинного академического труда.

Функциональный и дисфункциональный кентавр

Однако не всякое взаимодействие с КомИИ порождает ситуацию Кентавра.

Большинство актуальных практик использования языковых моделей в академической среде не являются ситуацией Кентавра. Отсюда вытекает вопрос, который нельзя обойти: рефлексиируют ли исследователи, работающие в ситуации Кентавра, эти принципы в одинаковой мере? Из сказанного следует необходимость явного формулирования критериев функционального кентавра в противовес дисфункциональному.

Функциональный кентавр – это режим взаимодействия с КомИИ, при котором все пять принципов реализуются совместно: итеративный диалог, в котором каждый ответ КомИИ конституирует следующий вопрос; трансформация исследовательской интенции в ходе работы; активная эпистемическая бдительность как процессуальная, а не разовая позиция; неделимость и принятие ответственности за каждое содержательное решение; эпистемическая локальность как устойчивость к незаметному концептуальному замещению. Дисфункциональный кентавр – это режим, при котором один или несколько из этих принципов нарушены: КомИИ используется для оформления уже готового содержания без подлинного диалога; или только для подтверждения готовых гипотез при игнорировании возражений; или его результаты принимаются без верификации; или следы взаимодействия скрываются, что означает уклонение от публичной ответственности.

Дисфункциональный кентавр – это кентавр, работающий неправильно. Задача состоит не в разоблачении имитации, а в диагностике нарушения и восстановлении функции. Традиционные инструменты детектирования ИИ-текста принципиально неспособны ответить на вопрос о функциональности Кентавра, поскольку они определяют, написан ли текст моделью, а не человеком, тогда как функциональный кентавр оставляет паттерны Автора, а не скриптора. Вместо детекции – документирование диалогов, обоснование решений, публичная рефлексия о характере взаимодействия с КомИИ. Ситуа-

ция Кентавра не снижает требований к исследователю – она их повышает.

Возрождение Автора: эпистемический кентавр против постструктуралистского скриптора

Принцип эмерджентности смысла вступает в прямой диалог с постструктуралистской концепцией смерти автора. В 1967 году Р. Барт провозгласил смерть автора. Суверенный субъект, чьи намерения определяют смысл текста, упраздняется [34]; на его место приходит «скриптор» – не творец, а переписчик, смешивающий коды культуры. Двумя годами позже М. Фуко уточнил тезис через понятие «функции-автора». Автор – дискурсивная функция, возникающая там, где текст – это то, что дозволено в конкретную эпоху [35].

Многие исследователи усматривают в языковых моделях материальное воплощение бартовского скриптора. ИИ смешивает коды культуры, не имея биографической личности и намерений. Ряд авторов утверждает, что генеративный ИИ знаменует «смерть функции-автора» [36]. Прежде чем возразить против этого диагноза, необходимо рассмотреть его наиболее сильную версию.

Наиболее серьёзное возражение против нашего тезиса о «возрождении Автора» состоит в следующем: итеративный диалог с КомИИ может быть описан именно как бартовская интертекстуальная игра, в которой «субъект» рассеивается не в меньшей степени, чем у традиционного скриптора. Ведь вопросы, которые исследователь задаёт КомИИ, сами формируются в пространстве уже существующих дискурсов, концептуальных привычек и языковых кодов и в этом смысле «навигатор диалога» столь же не является суверенным субъектом, как и переписчик. Этот аргумент достаточно силён, чтобы потребовать прямого ответа.

Наш ответ состоит в следующем. Р. Барт писал о смерти автора как структурном тезисе о способе производства смысла в эпоху, когда письмо мыслилось как воспроизведе-

ние предсуществующего означаемого. Тезис был направлен против идеологии творца, чьё биографическое намерение определяет единственно верный смысл текста. Но, даже принимая этот структурный тезис, мы обнаруживаем принципиальное различие между скриптором и навигатором диалога. Скриптор движется по уже проложенным путям дискурса; навигатор создаёт конкретный путь выбора в пространстве возможностей – выбор, который не предопределён никакой предшествующей структурой и не воспроизводим никаким другим навигатором. Именно уникальность этого пути, а не претензия на суверенитет субъекта конституирует авторство кентавра. Речь идёт не о восстановлении мёртвого Автора-творца, но о рождении нового типа авторства: авторства как ответственного выбора пути в диалоге.

Фукианский анализ раскрывает иное, более убедительное измерение. Современные исследователи выявляют проблему: кто несёт функцию-автора в тексте, произведённом при участии языковой модели? [37; 38]. Принцип неделимости ответственности даёт прямой ответ: именно потому, что КомИИ не обладает нормативной субъектностью и не может быть привлечён к ответственности, функция-автор не исчезает – она концентрируется в человеке. Чем активнее КомИИ участвует в производстве текста, тем острее вопрос о человеческой ответственности за него. Происходит не растворение, а усиление функции-автора. Смерть автора, провозглашённая Р. Бартом применительно к традиционному письму, оборачивается в ситуации Кентавра его возрождением в диалогическом облике.

Концепция эпистемического кентавра отличается от конкурирующих подходов, прежде всего от «дистанционного письма» Л. Флориды [38], в принципиальном отношении. «Дистанционное письмо» сохраняет иерархическую модель – человек-режиссёр управляет ИИ-исполнителем. Кентавр описывает диалогическое сопроизводство смысла, в котором интенция исследователя сама

трансформируется в процессе взаимодействия. Это не «дистанционное», а предельно близкое письмо.

Импликации для науки и высшего образования

Эпистемологические следствия. Концепция кентавра описывает новый режим производства знания с тремя специфическими характеристиками. Во-первых, невозпроизводимость пути при воспроизводимости результата. Конкретный диалог уникален, но результат и его обоснование верифицируемы [39]. Во-вторых, изменение темпоральности научного мышления. Ускорение фаз концептуализации создаёт риск недостаточной укоренённости в понимании [40]. В-третьих, алгоритмическая власть над исследовательской повесткой. КомИИ нормализует «усреднённые» вопросы. Функциональный кентавр использует «сбои» модели как диагностический сигнал о границах устоявшегося знания.

Авторство и атрибуция: принцип раскрытия. Позиции ведущих издательств запрещают указывать ИИ соавтором [41]. При всей практической необходимости они опираются на неявное допущение о том, что участие КомИИ аналогично участию человека-соавтора. Концепция кентавра предлагает иную рамку: КомИИ не является и не может быть соавтором, но его участие реально и должно быть раскрыто не как соавторство, но как методологическая информация. Мы утверждаем принцип раскрытия, а не атрибуции [37; 42].

Эпистемические риски культурной асимметрии. Концепция кентавра была бы неполной без анализа политического измерения познания – политического в эпистемологическом смысле. Эпистемическая ненейтральность КомИИ, сформулированная в пятом принципе (см. раздел «Пять принципов эпистемического кентавра»), приобретает в контексте академического производства знания задокументированные и измеримые последствия. Эта асимметрия зафиксирована в ра-

стущем корпусе критических исследований; языковые технологии по своему устройству благоприятствуют определённым языкам и диалектам, что ведёт к системам, ограниченным в выражении социокультурно значимых понятий незападных сообществ [43]. Языковые модели рискуют воспроизводить специфическую форму несправедливости, которую М. Фрикер назвала эпистемической несправедливостью [44]. Западные традиции мысли, западные социальные реалии, в том числе богатейшая русскоязычная гуманитарная традиция, системно недопредставлены в обучении моделей, что создаёт риск «эпистемцида» – незаметного вытеснения альтернативных способов познания [45].

Теоретическое основание для осмысления этой проблемы предоставляет концепция «ситуированных знаний» Д. Харауэй [46]: не существует «взгляда из ниоткуда», всякое знание производится из конкретной позиции [47]. Д. Харауэй тем самым предоставляет философское основание пятому принципу: «нейтральность» КомИИ – это частная перспектива, утратившая видимость в силу доминирования породившего её дискурса.

Принцип эпистемической локальности (см. раздел «Пять принципов эпистемического кентавра») обнаруживает здесь своё операциональное значение. В критическом измерении он предписывает распознавать латуровские «переводы» в действии: когда модель встраивает, например, концепцию зоны ближайшего развития в схемы западной педагогической психологии, изымая её из деятельностного контекста, происходит именно такое концептуальное замещение. В продуктивном измерении тот же принцип превращает риск эпистемцида в исследовательский ресурс: деятельностная психология и семиосфера Ю.М. Лотмана вводятся в диалог не как «региональные варианты», а как концептуальные языки, описывающие то, для чего у западной науки нет готовых категорий.

Эпистемическая локальность органично связана с принципом неделимости ответ-

ственности. Если ответственность за смысл принадлежит человеку, то и ответственность за концептуальную рамку тоже. Делегировать выбор концептуального языка модели – значит делегировать ей не рутину, а исходные допущения исследования. В этом смысле эпистемическая бдительность должна включать бдительность культурную и концептуальную: постоянную готовность спросить, чьими категориями мыслит в данный момент кентавр – своими или навязанными.

Компетенции познания. Эпоха кентавра добавляет к традиционным исследовательским компетенциям четвертую область – компетенции познания. Они включают: искусство вопрошания (формулировать вопросы, открывающие новые горизонты); критическую рецепцию генеративного текста (читать ответы как гипотезы); интеграцию диалога в понимание (ассимилировать результаты настолько, чтобы защищать их и опровергать); методологическую рефлексию (документировать характер взаимодействия).

Снижение требований к предметной подготовке под предлогом доступности КомИИ – прямой путь к производству дисфункциональных кентавров. Кентавр усиливает сильных и делает слабость слабых видимой.

Сообщество практик кентавра. Ни индивидуальная эпистемическая бдительность, ни институциональные регламенты не являются достаточным ответом на вызовы, которые ставит перед академическим сообществом эпоха кентавра. Регламенты описывают нормы, но не воспроизводят практику; индивидуальная этика удерживает стандарт лишь до тех пор, пока есть среда, которая его разделяет и верифицирует. Между регламентом и индивидуальной совестью – пустота, которую может заполнить только профессиональное сообщество.

Для концептуализации этой роли мы опираемся на понятие «сообществ практик будущего», введенное Д.И. Земцовым применительно к сообществам, которые способ-

ны обеспечить социокультурное освоение технологических изменений [48]. В отличие от «сообществ практик» Э. Венгера [49], ориентированных на воспроизводство сложившихся профессиональных норм внутри устойчивого социального поля, сообщество практик будущего принципиально характеризуется невключённостью в такое поле: оно не встроено в академическую иерархию, корпоративную структуру или государственный институт, но при этом встроено в общество через заботу об общественном благе. Именно эта позиция вне поля при одновременной включённости в общество делает сообщество практик будущего агентом освоения, а не просто агентом воспроизводства. Ключевым механизмом его влияния является транзиторность: участники уходят, но переносят отрефлексированные практики во внешние социальные поля, сохраняя тем самым преемственность норм при смене конкретных носителей. Применительно к ситуации Кентавра это означает: формирование функциональных практик познания невозможно усилиями одиночек и неустойчиво в рамках официальных институций – оно требует именно того типа горизонтальных сообществ, который описывает Д.И. Земцов: инициативных, инновационных, неформальных и ориентированных на передачу практик, а не только знаний.

Мы призываем к формированию *сообщества практик кентавра* – горизонтальной сети исследователей, разделяющих ценности кентаврического познания. Сообщество практик кентавра отличается от существующих академических институций в одном принципиальном отношении: его центральным объектом является не знание как продукт, но *процесс его производства* – итеративный, диалогический, эмерджентный. Это означает культуру, в которой делиться диалогами столь же естественно, как делиться данными; где рефлексия о методе не является приложением к статье, но конституирует саму исследовательскую идентичность. Именно такую культуру академическое со-

общество умело строить применительно к методам исследования на протяжении десятилетий и именно её предстоит выстроить заново применительно к кентаврическому познанию.

Заключение: искусственный интеллект обостряет человеческое в человеке

Как мы показали, КомИИ не удовлетворяет критериям расширенного разума. Смежные теории схватывают важное, но упускают принципиальное – семантическую пластичность и неделимость ответственности. Пять принципов кентавра образуют единую архитектуру от онтологии к этике. Смерть автора оборачивается его возрождением. Функциональный кентавр отличим от дисфункционального.

Зеркало не создаёт отражение – оно его обнажает. КомИИ не создаёт интеллект исследователя и не заменяет его. Он делает видимым то, что было в нём всегда, но скрывалось за рутиной академического производства: способность или неспособность к подлинному вопрошанию, к удержанию интеллектуального напряжения, к принятию ответственности за мысль. Это не жестокость КомИИ – это его честность. И в этой честности состоит финальный в нашей статье парадокс эпохи кентавра: искусственный интеллект обостряет человеческое в человеке.

Этот тезис может быть прочитан двояко. В первом прочтении – это тезис об усилении. КомИИ поднимает на порядок когнитивные возможности того, кто способен работать в режиме функционального кентавра. Человеческий интеллект в функциональном единстве с КомИИ становится мощнее именно потому, что остаётся человеческим: интенциональным, ответственным, способным к подлинному вопрошанию. Во втором прочтении – это тезис об обнажении рисков. КомИИ обостряет угрозы интеллектуальной деградации и эпистемической трусости. Оба прочтения сходятся в том, что КомИИ – усилитель того, что уже есть в человеке.

Именно те качества, которых КомИИ не имеет и иметь не может, – способность задавать вопросы из подлинного незнания, выдерживать интеллектуальное напряжение неопределённости, принимать ответственность за мысль – суть то, что делает человека неустранимым субъектом познания. Широкое распространение КомИИ поставило непростые вопросы о фундаментальных свойствах человека. И оказалось, что главный вопрос при обсуждении искусственного интеллекта по-прежнему о человеке.

Литература

1. Sparrow B., Liu J., Wegner D.M. Google effects on memory // *Science*. – 2011. – Vol. 333. – P. 776–778. – DOI: 10.1126/science.1207745.
2. Ward A.F. Supernormal: how the internet is changing our memories and our minds // *Psychological Inquiry*. – 2013. – Vol. 24. – P. 341–348. – DOI: 10.1080/1047840X.2013.850148.
3. Messeri L., Crockett M.J. Artificial intelligence and illusions of understanding in scientific research // *Nature*. – 2024. – Vol. 627, no. 8002. – P. 49–58. – DOI: 10.1038/s41586-024-07146-0.
4. Clark A. Extending minds with generative AI // *Nature Communications*. – 2025. – Vol. 16. – Article no. 4627. – DOI: 10.1038/s41467-025-59906-9.
5. Clark A., Chalmers D. The extended mind // *Analysis*. – 1998. – Vol. 58, no. 1. – P. 7–19. – DOI: 10.1093/analys/58.1.7.
6. Никольский В.С. Коммуникативный искусственный интеллект: концептуализация новой реальности в образовании // *Высшее образование в России*. – 2025. – Т. 34, № 6. – С. 152–168. – DOI: 10.31992/0869-3617-2025-34-6-152-168.
7. Sison A.J.G., Daza M.T., Gozalo-Brizuela R., Garrido-Merchán E.C. ChatGPT: more than a “weapon of mass deception” // *International Journal of Human-Computer Interaction*. – 2023. – Vol. 40, no. 17. – P. 1–20. – DOI: 10.1080/10447318.2023.2225931.
8. Hutchins E. *Cognition in the Wild*. – MIT Press, 1995. – 381 p. – ISBN: 978-0-262-08231-0.
9. Grinschgl S., Neubauer A.C. Supporting cognition with modern technology: distributed cognition today and in an AI-enhanced future // *Frontiers in Artificial Intelligence*. – 2022. – Vol. 5. – Article no. 908261. – DOI: 10.3389/frai.2022.908261.
10. Латур Б. Наука в действии: следуя за учёными и инженерами внутри общества / Б. Латур ; перевод с французского К. Федоровой. – Санкт-Петербург : Изд-во Европейского университета в Санкт-Петербурге, 2013. – 414 с. – ISBN: 978-5-94380-161-7.
11. Venturini T. Bruno Latour and Artificial Intelligence // *Tecnoscienza – Italian Journal of Science & Technology Studies*. – 2023. – Vol. 14, no. 2. – P. 101–114. – DOI: 10.6092/issn.2038-3460/18359.
12. Gutiérrez J.L.M. On actor-network theory and algorithms: ChatGPT and the new power relationships in the age of AI // *AI and Ethics*. – 2024. – Vol. 4. – P. 1071–1084. – DOI: 10.1007/s43681-023-00314-4.
13. Dean A., Orfanoudaki A., Saghaian S., Song K., Chakker A.H.A., Cook C. Algorithm, Human, or the Centaur: How to Enhance Clinical Care? // *SSRN Electronic Journal*. – 2022. – DOI: 10.2139/ssrn.4302002.
14. Saghaian S., Idan L. Effective Generative AI: The Human-Algorithm Centaur // *Harvard Data Science Review*. – 2024. – DOI: 10.1162/99608f92.19d78478.
15. Haupt A. Position: AI Should Not Be An Imitation Game: Centaur Evaluations // *Stanford Digital Economy Lab Working Paper*. – 2025. – URL: <https://digitaleconomy.stanford.edu/publications/centaur-evaluations/> (дата обращения: 30.04.2026).
16. Солдатова Г.У., Чигарькова С.В., Илюхина С.Н. Метаморфозы идентичности человека достроенного: от цифрового донора к цифровому кентавру // *Социальная психология и общество*. – 2024. – Т. 15, № 3. – С. 23–43. – DOI: 10.17759/sps.2024150302.
17. Soldatova G.U., Chigarkova S.V., Ilyukhina S.N. The Digital Centaur as a Type of Technologically Augmented Human in the AI Era: Personal and Digital Predictors // *Behavioral Sciences*. – 2025. – Vol. 15, no. 3. – Article no. 289. – DOI: 10.3390/bs15111487.
18. Гаранин М.А., Максименко А.Ю., Веляева К.С., Золотарева В.В. Цифровые «кентавры» в образовании // *Экономика, предпринимательство и право*. – 2024. – Т. 14, № 11. – С. 5363–5382. – DOI: 10.18334/epp.14.11.122032.
19. Bratman M. *Intention, Plans, and Practical Reason*. – Harvard University Press, 1987. – 224 p. – ISBN: 978-0-674-45818-5.

20. Guzman A.L., Lewis S.C. Artificial intelligence and communication: A Human-Machine Communication research agenda // *New Media & Society*. – 2020. – Vol. 22, no. 1. – P. 70–86. – DOI: 10.1177/1461444819858691.
21. Coeckelbergh M., Gunkel D. *Communicative AI: A Critical Introduction to Large Language Models*. – Polity Press, 2024. – 144 p. – ISBN: 978-1-5095-6759-1.
22. Chalmers D.J. ‘Strong and Weak Emergence’ / Philip Clayton, and Paul Davies (eds). *The Re-Emergence of Emergence: The Emergentist Hypothesis from Science to Religion*. – Oxford, 2008; online edn, Oxford Academic, 3 Oct. 2011. – DOI: 10.1093/acprofoso/9780199544318.003.0011.
23. Holland J.H. *Emergence: From Chaos to Order*. – Oxford University Press, 1998. – 258 p. – ISBN: 978-0-19-850409-2.
24. Платон. Тезтет / пер. Т.В. Васильевой // Платон. Собрание сочинений: в 4 т. – Москва: Мысль, 1993. – Т. 2. – С. 192–275. – ISBN: 5-244-00471-9.
25. Выготский А.С. Мышление и речь. – Москва: АСТ, 2025. – 576 с. – ISBN: 9785171771683.
26. Бахтин М.М. Вопросы литературы и эстетики: Исследования разных лет. – Москва : Худож. лит., 1975. – 502 с. – URL: https://biblio.imli.ru/images/abook/teoriya/Bahtin_M.M._Voprosy_literaturny_i_estetiki_1975.pdf (дата обращения: 30.04.2026).
27. Austin J.L. *How to Do Things with Words*. – Oxford: Oxford University Press, 1962. – 178 p. – ISBN: 978-0-674-41152-4.
28. Searle J.R. *Speech Acts: An Essay in the Philosophy of Language*. – Cambridge: Cambridge University Press, 1969. – 203 p. – ISBN: 978-0-521-09626-3.
29. Жинкин Н.И. О кодовых переходах во внутренней речи // *Вопросы языкознания*. – 1964. – № 6. – С. 26–38.
30. Жинкин Н.И. *Речь как проводник информации*. – Москва: Наука, 1982. – 157 с.
31. Shanahan M. *Talking About Large Language Models*. – arXiv:2212.03551, 2023. – DOI: 10.48550/arXiv.2212.03551.
32. Jonas H. *The Imperative of Responsibility: In Search of an Ethics for the Technological Age*. – University of Chicago Press, 1984. – 255 p. – ISBN: 978-0-226-40596-4.
33. Jasanoff S. *Designs on Nature: Science and Democracy in Europe and the United States*. – Princeton University Press, 2005. – 392 p. – ISBN: 978-0-691-11811-6.
34. Барт Р. Смерть автора // *Избранные работы: Семиотика. Поэтика*. – Москва: Прогресс, 1989. – С. 384–391.
35. Фуко М. Что такое автор? // *Воля к истине: по ту сторону знания, власти и сексуальности*. – Москва: Касталь, 1996. – С. 7–46. – ISBN: 5-85374-006-7.
36. Slater A. Phantoms of citation: AI and the death of the author-function // *Poetics Today*. – 2024. – Vol. 45, no. 2. – P. 223–248. – DOI: 10.1215/0335372-11092818.
37. Gretzky M., Dishon G. Algorithmic-authors in academia: blurring the boundaries of human and machine knowledge production // *Learning, Media and Technology*. – 2025. – DOI: 10.1080/17439884.2025.2452196.
38. Floridi L. Distant Writing: Literary Production in the Age of Artificial Intelligence. *Minds & Machines*. – 2025. – Vol. 35, no. 3. – DOI: 10.1007/s11023-025-09732-1.
39. Flick U. *An Introduction to Qualitative Research*. – 6th ed. Sage, 2018. – 696 p. – ISBN: 978-1-5264-4565-0.
40. Ou A.W., Stöhr C., Malmström H. Academic communication with AI-powered language tools in higher education: From a post-humanist perspective // *System*. – 2024. – Vol. 121. – DOI: 10.1016/j.system.2024.103225.
41. COPE Council. *Authorship and AI tools*. – Committee on Publication Ethics, 2023. – DOI: 10.24318/cvrbms.
42. Ананин Д.П., Комаров Р.В., Реморенко И.М. «Когда честно – хорошо, для имитации – плохо»: стратегии использования генеративного искусственного интеллекта в российском вузе // *Высшее образование в России*. – 2025. – Т. 34, № 2. – С. 31–50. – DOI: 10.31992/0869-3617-2025-34-2-31-50.
43. Helm P., Bella G., Koch G., Giunchiglia F. Diversity and language technology: how language modeling bias causes epistemic injustice // *Ethics and Information Technology*. – 2024. – DOI: 10.1007/s10676-023-09742-6.
44. Kraft A., Soulier E. Knowledge-Enhanced Language Models Are Not Bias-Proof: Situated Knowledge and Epistemic Injustice in AI // *Proceedings of FAccT '24*. – ACM, 2024. – DOI: 10.1145/3630106.3658981.
45. Olaniyan Y.D., Martins M.O., Al Maqrashi R.H. Generative artificial intelligence and epistemic (in)justice: perspectives from higher education students in the Global North and South

- // *Frontiers in Human Dynamics*. – 2026. – Vol. 8. – Article no. 1790324. – DOI: 10.3389/fhumd.2026.1790324.
46. Haraway D. Situated Knowledges: The Science Question in Feminism and the Privilege of Partial Perspective // *Feminist Studies*. – 1988. – Vol. 14, no. 3. – P. 575–599. – DOI: 10.2307/3178066.
47. Trächtler J. The world as witty agent – Donna Haraway on the object of knowledge // *Frontiers in Psychology*. – 2024. – DOI: 10.3389/fpsyg.2024.1389575.
48. Земцов Д.И. Сообщества практик будущего в российских университетах: фаблабы, ЦМИ-Ты, кружки // *Высшее образование в России*. – 2023. – Т. 32, № 5. – С. 36–55. – DOI: 10.31992/0869-3617-2023-32-5-36-55.
49. Wenger E. *Communities of Practice: Learning, Meaning, and Identity*. – Cambridge University Press, 1998. – 318 p. – ISBN: 978-0-521-43017-3.

Статья поступила в редакцию 11.05.2026

Принята к публикации 16.06.2026

References

1. Sparrow, B., Liu, J., Wegner, D.M. (2011). Google Effects on Memory. *Science*. Vol. 333, pp. 776–778, doi: 10.1126/science.1207745.
2. Ward, A.F. (2013). Supernormal: How the Internet Is Changing Our Memories and Our Minds. *Psychological Inquiry*. Vol. 24, pp. 341–348, doi: 10.1080/1047840X.2013.850148.
3. Messeri, L., Crockett, M.J. (2024). Artificial Intelligence and Illusions of Understanding in Scientific Research. *Nature*. Vol. 627, no. 8002, pp. 49–58, doi: 10.1038/s41586-024-07146-0.
4. Clark, A. (2025). Extending Minds with Generative AI. *Nature Communications*. Vol. 16, article no. 4627, doi: 10.1038/s41467-025-59906-9.
5. Clark, A., Chalmers, D. (1998). The Extended Mind. *Analysis*. Vol. 58, no. 1, pp. 7–19, doi: 10.1093/analys/58.1.7.
6. Nikolsky, V.S. (2025). Communicative Artificial Intelligence: Conceptualization of a New Reality in Education. *Vysshee Obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 152–168, doi: 10.31992/0869-3617-2025-34-6-152-168 (In Russ., abstract in Eng.).
7. Sison, A.J.G., Daza, M.T., Gozalo-Brizuela, R., Garrido-Merchán, E.C. (2023). ChatGPT: More Than a “Weapon of Mass Deception”. *International Journal of Human-Computer Interaction*. Vol. 40, no. 17, pp. 1–20, doi: 10.1080/10447318.2023.2225931.
8. Hutchins, E. (1995). *Cognition in the Wild*. MIT Press. 381 p. ISBN: 978-0-262-08231-0.
9. Grinschgl, S., Neubauer, A.C. (2022). Supporting Cognition with Modern Technology: Distributed Cognition Today and in an AI-Enhanced Future. *Frontiers in Artificial Intelligence*. Vol. 5, article no. 908261, doi: 10.3389/frai.2022.908261.
10. Latour, B. (1987). *Science in Action: How to Follow Scientists and Engineers Through Society*. Harvard University Press (Russian Translation by K. Fedorova, Saint Oetersburg: Izdatel'stvo Evropeyskogo Universiteta v Sankt-Peterburge, 2013, 414 p.).
11. Venturini, T. (2023). Bruno Latour and Artificial Intelligence. *Tecnoscienza – Italian Journal of Science & Technology Studies*. Vol. 14, no. 2, pp. 101–114, doi: 10.6092/issn.2038-3460/18359.
12. Gutiérrez, J.L.M. (2024). On Actor-Network Theory and Algorithms: ChatGPT and the New Power Relationships in the Age of AI. *AI and Ethics*. Vol. 4, pp. 1071–1084, doi: 10.1007/s43681-023-00314-4.
13. Dean A., Orfanoudaki A., Saghaian S., Song K., Chakkera H.A., Cook C. (2022). Algorithm, Human, or the Centaur: How to Enhance Clinical Care? *SSRN Electronic Journal*. Doi: 10.2139/ssrn.4302002.
14. Saghaian, S., Idan, L. (2024). Effective Generative AI: The Human-Algorithm Centaur. *Harvard Data Science Review*. Doi: 10.1162/99608f92.19d78478.

15. Haupt, A. (2025). Position: AI Should Not Be an Imitation Game: Centaur Evaluations. *Stanford Digital Economy Lab Working Paper*. Stanford University. Available at: <https://digitaleconomy.stanford.edu/publications/centaur-evaluations/> (accessed 30.04.2026).
16. Soldatova, G.U., Chigar'kova, S.V., Ilyukhina, S. N. (2024). Metamorphoses of the Identity of an Augmented Human: From a Digital Donor to a Digital Centaur. *Sotsial'naya Psikhologiya i Obschestvo*. Vol. 15, no. 3, pp. 23-43, doi: 10.17759/sps.2024150302 (In Russ., abstract in Eng.).
17. Soldatova, G.U., Chigarkova, S.V., Ilyukhina, S.N. (2025). The Digital Centaur as a Type of Technologically Augmented Human in the AI Era: Personal and Digital Predictors. *Behavioral Sciences*. Vol. 15, no. 3, article no. 289, doi: 10.3390/bs15030289 (In Russ., abstract in Eng.).
18. Garanin, M.A., Maksimenko, A.Yu., Velyaeva, K.S., Zolotareva, V.V. (2024). Digital "Centaur" in Education. *Ekonomika, predprinimatel'stvo i pravo = Journal of Economics, Entrepreneurship and Law*. Vol. 14, no. 11, pp. 5363-5382, doi: 10.18334/epp.14.11.122032 (In Russ., abstract in Eng.).
19. Bratman, M. (1987). *Intention, Plans, and Practical Reason*. Harvard University Press. 224 p. ISBN: 978-0-674-45818-5.
20. Guzman, A.L., Lewis, S.C. (2020). Artificial Intelligence and Communication: A Human-Machine Communication Research Agenda. *New Media & Society*. Vol. 22, no. 1, pp. 70-86, doi: 10.1177/1461444819858691.
21. Coeckelbergh, M., Gunkel, D. (2024). Communicative AI: A Critical Introduction to Large Language Models. *Polity Press*. 144 p. ISBN: 978-1-5095-6759-1.
22. Chalmers, D.J. (2008). 'Strong and Weak Emergence'. In: Philip Clayton, and Paul Davies (eds). *The Re-Emergence of Emergence: The Emergentist Hypothesis from Science to Religion*. Oxford, 2008; online edn, Oxford Academic, 3 Oct. 2011, doi: 10.1093/acprof:oso/9780199544318.003.0011.
23. Holland, J.H. (1998). *Emergence: From Chaos to Order*. Oxford University Press. 258 p. ISBN: 978-0-19-850409-2.
24. Plato. (1993). *Theaetetus* (T.V. Vasilieva, Trans.). In: *Platon. Sobranie sochineniy v 4 tomakh* (Vol. 2, pp. 192-275). Mysl [Thought]. ISBN: 5-244-00471-9. (In Russ.).
25. Vygotsky, L.S. (2025). *Myslenie i rech'* [Thinking and Speech]. AST. 576 p. ISBN: 9785171771683. (In Russ.).
26. Bakhtin, M.M. (1975). *Voprosy literatury i estetiki: Issledovaniya raznykh let* [Questions of Literature and Aesthetics: Researches of Different Years]. Khudozhestvennaya Literatura. 502 p. Available at: https://biblio.imli.ru/images/abook/teoriya/Bakhtin_M.M._Voprosy_literatury_i_estetiki_1975.pdf (accessed 30.04.2026). (In Russ.).
27. Austin, J.L. (1962). *How to Do Things with Words*. Oxford: Oxford University Press. 178 p. ISBN: 978-0-674-41152-4.
28. Searle, J.R. (1969). *Speech Acts: An Essay in the Philosophy of Language*. Cambridge University Press. 203 p. ISBN: 978-0-521-09626-3.
29. Zhinkin, N.I. (1964). O kodovykh perekhodakh vo vnutrenney rechi [On Code Transitions in Inner Speech]. *Voprosy Yazykoznaniiya* [Topics in the Study of Language]. Vol. 6, pp. 26-38. (In Russ.).
30. Zhinkin, N.I. (1982). *Rech' kak provodnik informatsii* [Speech as a Conductor of Information]. Moscow: Nauka [Science]. 157 p. (In Russ.).
31. Shanahan, M. (2023). Talking about Large Language Models. *arXiv:2212.03551*. Doi: 10.48550/arXiv.2212.03551.
32. Jonas, H. (1984). *The Imperative of Responsibility: In Search of an Ethics for the Technological Age*. University of Chicago Press. 255 p. ISBN: 978-0-226-40596-4.
33. Jasanoff, S. (2005). *Designs on Nature: Science and Democracy in Europe and the United States*. Princeton University Press. 392 p. ISBN: 978-0-691-11811-6.

34. Barthes, R. (1989). Smert' avtora [The Death of the Author]. In: *Izbrannye raboty: Semiotika. Poetika*. Pp. 384-391. Progress [Progress]. (In Russ.).
35. Foucault, M. (1996). Chto takoe avtor? [What Is an Author?]. In: *Volya k istine: po tu storonu znaniya, vlasti i seksual'nosti*. Pp. 7-46. Kastal'. ISBN: 5-85374-006-7. (In Russ.).
36. Slater, A. (2024). Phantoms of Citation: AI and the Death of the Author-Function. *Poetics Today*. Vol. 45, no. 2, pp. 223-248, doi: 10.1215/03335372-11092818.
37. Gretzky, M., Dishon, G. (2025). Algorithmic-Authors in Academia: Blurring the Boundaries of Human and Machine Knowledge Production. *Learning, Media and Technology*. Doi: 10.1080/17439884.2025.2452196.
38. Floridi, L. (2025) Distant Writing: Literary Production in the Age of Artificial Intelligence. *Minds & Machines*. Vol. 35, no. 3, doi: 10.1007/s11023-025-09732-1.
39. Flick, U. (2018). *An Introduction to Qualitative Research* (6th ed.). Sage. 696 p. ISBN: 978-1-5264-4565-0.
40. Ou, A. W., Stöhr, C., Malmström, H. (2024). Academic Communication with AI-Powered Language Tools in Higher Education: From a Post-Humanist Perspective. *System*. Vol. 121, article no. 103225, doi: 10.1016/j.system.2024.103225.
41. COPE Council. (2023). *Authorship and AI Tools*. Committee on Publication Ethics. Doi: 10.24318/cCVRZBms.
42. Ananin, D.P., Komarov, R.V., Remorenko, I.M. (2025). "When Honest It's Good, for Imitation It's Bad": Strategies for Using Generative Artificial Intelligence in a Russian University. *Vysshee Obrazovanie v Rossii = Higher education in Russia*. Vol. 34, no. 2, pp. 31-50, doi: 10.31992/0869-3617-2025-34-2-31-50 (In Russ., abstract in Eng.).
43. Helm, P., Bella, G., Koch, G., Giunchiglia, F. (2024). Diversity and Language Technology: How Language Modeling Bias Causes Epistemic Injustice. *Ethics and Information Technology*. Doi: 10.1007/s10676-023-09742-6.
44. Kraft, A., Soulier, E. (2024). Knowledge-Enhanced Language Models Are Not Bias-Proof: Situated Knowledge and Epistemic Injustice in AI. In: *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency (FAccT'24)*. ACM. Doi: 10.1145/3630106.3658981.
45. Olaniyan, Y.D., Martins, M.O., Al Maqrashi, R.H. (2026). Generative Artificial Intelligence and Epistemic (In)Justice: Perspectives from Higher Education Students in the Global North and South. *Frontiers in Human Dynamics*. Vol. 8, article no. 1790324, doi: 10.3389/fhumd.2026.1790324.
46. Haraway, D. (1988). Situated Knowledges: The Science Question in Feminism and the Privilege of Partial Perspective. *Feminist Studies*. Vol. 14, no. 3, pp. 575-599, doi: 10.2307/3178066.
47. Trächtler, J. (2024). The World as Witty Agent – Donna Haraway on the Object of Knowledge. *Frontiers in Psychology*. Doi: 10.3389/fpsyg.2024.1389575.
48. Zemtsov, D.I. (2023). Communities of Practice of the Future in Russian Universities: Fablabs, Centers for Youth Innovative Creativity, Kruzhoks. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 5, pp. 36-55, doi: 10.31992/0869-3617-2023-32-5-36-55 (In Russ., abstract in Eng.).
49. Wenger, E. (1998). *Communities of Practice: Learning, Meaning, and Identity*. Cambridge University Press. 318 p. ISBN: 978-0-521-43017-3.

The paper was submitted 11.05.2026

Accepted for publication 16.06.2026

От разрешительных практик к институциональным нормам: этическое регулирование использования искусственного интеллекта в российском высшем образовании

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-31-54

Филипова Александра Геннадьевна – д-р социол. наук, зав. лабораторией комплексных исследований детства¹, старший научный сотрудник Института детства², Scopus Author ID: 57190731212, ORCID: 0000-0002-7475-1961, Alexgen77@list.ru

¹ Владивостокский государственный университет, Владивосток, Россия

Адрес: 690014, г. Владивосток, ул. Гоголя, д. 41

² Российский государственный педагогический университет им. А.И. Герцена, Санкт-Петербург, Россия

Адрес: 191186, г. Санкт-Петербург, Набережная реки Мойки, д. 48

Склярва Софья Андреевна – ассистент кафедры международных отношений и государственного управления, SPIN-код: 8811-9726, sofya.sklyarova.96@mail.ru

Владивостокский государственный университет, Владивосток, Россия

Адрес: 690014, г. Владивосток, ул. Гоголя, д. 41

***Аннотация.** Актуальность исследования обусловлена противоречием между широким фактическим использованием искусственного интеллекта (ИИ) в российских университетах и дефицитом формализованных правил, регулирующих этот процесс. Целью работы является комплексный анализ официальных документов российских вузов, регламентирующих использование искусственного интеллекта в образовательном процессе, в т. ч. этические аспекты и риски. Методология базируется на двухэтапном контент-анализе открытых нормативных документов университетов. Теоретическая рамка исследования опирается на положения неoinституциональной теории, включая концепции институционального изоморфизма и стадий институционализации. Научная новизна исследования состоит не только в географическом расширении поля исследований политики ИИ за счёт российского кейса, но и в развитии неoinституционального анализа регуляций в условиях технологической неопределённости. Российский материал позволяет уточнить динамику ранней институционализации: показано, что формирование ИИ-политик носит нелинейный характер, сочетающий заимствование дискурсивных формул и сохранение «институционального вакуума» на уровне процедур. Вводится типология адаптивного и ограничительного режимов регуляции, интерпретируемых как два устойчивых режима институционального реагирования на цифровые инновации, и демонстрируется их одновременное сосуществование в одной национальной системе. Этические принципы трактуются как специфический механизм символической легитимации*

вуза в ситуации дефицита детализированных регламентов, что дополняет существующие подходы к анализу институционального изоморфизма и конкурентной дифференциации университетов. Перспективы исследования рассматриваются в направлении анализа процессов интернализации правил участниками образовательного процесса и лонгитюдного наблюдения за эволюцией институциональных форм.

Ключевые слова: искусственный интеллект, генеративный искусственный интеллект, этические принципы, институционализация, легитимация, высшее образование, контент-анализ, неoinституциональная теория, академическая честность, политика использования ИИ, институциональный изоморфизм

Для цитирования: Филипова А.Г., Складорова С.А. От разрешительных практик к институциональным нормам: этическое регулирование использования искусственного интеллекта в российском высшем образовании // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 31–54. – DOI: 10.31992/0869-3617-2026-35-6-31-54.

From Permissive Practices to Institutional Norms: Ethical Regulation of Artificial Intelligence in Russian Higher Education

Original article

DOI: 10.31992/0869-3617-2026-35-6-31-54

Alexandra G. Filipova – Dr.Sci. (Sociology), Head of the Laboratory for Integrated Research on Childhood¹, Senior Researcher², Institute of Childhood, Scopus Author ID: 57190731212, ORCID: 0000-0002-7475-1961, Alexgen77@list.ru

¹ Vladivostok State University, Vladivostok, Russian Federation

Address: 41 Gogolya str., Vladivostok, 690014, Russian Federation

² Herzen State Pedagogical University of Russia, St. Petersburg, Russian Federation

Address: 48 Moika River emb., St. Petersburg, 191186, Russian Federation

Sofia A. Skliarova – Assistant at the Department of International Relations and Public Administration, SPIN-code: 8811-9726, sofya.sklyarova.96@mail.ru,

Vladivostok State University, Vladivostok, Russian Federation

Address: 41 Gogolya str., Vladivostok, 690014, Russian Federation

Abstract. The relevance of this study stems from a contradiction between the widespread de facto use of artificial intelligence (AI) in Russian universities and the lack of formal rules governing this process. The aim of this work is to comprehensively analyze official documents from Russian universities that regulate the use of AI in education, including ethical aspects and risks. The methodology is based on a two-stage content analysis of publicly available university policy documents. The theoretical framework draws on neo-institutional theory, particularly the concepts of institutional isomorphism and the stages of institutionalization. The scientific novelty of the study lies not only in expanding the geography of AI policy research by including the Russian case but also in advancing neo-institutional analysis of regulation under conditions of technological uncertainty. The Russian material allows us to refine the dynamics of early-stage institutionalization: we show that the formation of AI policies is nonlinear, combining the borrowing of discursive formulas with the persistence of an “institutional vacuum” at the procedural level. We introduce a typology of adaptive and

restrictive regulatory modes, interpreted as two stable patterns of institutional response to digital innovation, and demonstrate their simultaneous coexistence within a single national system. Ethical principles are interpreted as a specific mechanism for the symbolic legitimation of universities in the absence of detailed regulations. This complements existing approaches to analyzing institutional isomorphism and competitive differentiation among universities. Future research directions include analyzing how participants in the educational process internalize these rules and conducting longitudinal observations of the evolution of institutional forms.

Keywords: artificial intelligence, generative artificial intelligence, ethical principles, institutionalization, legitimation, higher education, content analysis, neo-institutional theory, academic integrity, AI policy, institutional isomorphism

Cite as: Filipova, A.G., Skliarova, S.A. (2026). From Permissive Practices to Institutional Norms: Ethical Regulation of Artificial Intelligence in Russian Higher Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 31-54, doi: 10.31992/0869-3617-2026-35-6-31-54 (In Russ., abstract in Eng.).

Введение

Использование технологий искусственного интеллекта (ИИ), в особенности генеративных моделей (ГИИ), в сфере высшего образования стало одним из наиболее значимых мировых трендов. Импульсом к этому масштабному процессу послужил публичный релиз платформы *ChatGPT* от компании *OpenAI* в ноябре 2022 года, который продемонстрировал беспрецедентную доступность и мощь ИИ-инструментов для широкой аудитории. Стремительное проникновение ГИИ в академическую среду было оценено экспертами и руководством вузов двояко: как серьезный технологический вызов, несущий риски для образования, и как новая возможность для трансформации обучения.

В ответ на эту двойственность университеты по всему миру начали активную работу по выработке системных подходов к интеграции ИИ. Эта работа выражается в разработке и внедрении локальных политик, призванных регулировать использование технологий ИИ в образовании.

С точки зрения неoinституциональной теории процесс выработки университетами общих политик в отношении ИИ может быть рассмотрен как проявление институционального изоморфизма — стремления организаций к единообразию для обретения

легитимности в условиях технологической неопределенности. Классическая типология институционального изоморфизма, разработанная П. ДиМаджио и У. Пауэллом, позволяет систематизировать силы, побуждающие вузы к этому: принудительное давление со стороны государства, миметическое подражание лидерам и нормативное влияние профессиональных стандартов [1]. Эта теоретическая рамка представляется значимой и для анализа текущей ситуации с регулированием ИИ в высшем образовании.

Неоинституциональный подход находит широкое применение в исследованиях внедрения инноваций в высшее образование. В австралийских университетах массовое внедрение онлайн-обучения в 1990-х годах происходило не в силу рационального выбора, а под действием механизмов институционального изоморфизма: принудительного давления государства, нормативного влияния профессиональных сообществ и миметического копирования успешных практик в условиях неопределенности [2]. Аналогичные механизмы эмпирически подтверждены при анализе внедрения аккредитационных процедур в Колумбии [3] и цифровизации образования в Малайзии [4]. Применительно к ИИ неoinституциональный подход позволяет рассматривать его не просто как технологию, а как формирующийся социальный

институт, структурирующий взаимодействия через формальные и неформальные правила и ставящий вопрос о его легитимности в академической среде [5]. Проникновение технологий ИИ в высшее образование выступает катализатором полного цикла институционализации, который включает возникновение повторяющихся практик, их объективацию в университетских политиках и последующее принятие новых норм академическим сообществом, согласно подходу П. Бергера и Т. Лукмана [6].

Российские университеты представляют собой уникальное поле для исследования, так как систематических исследований, направленных на содержательный анализ стратегической позиции вузов и качественных различий в их локальных политиках, явно недостаточно.

Целью данного исследования является комплексный анализ того, как российские вузы концептуализируют и регламентируют использование искусственного интеллекта в образовательном процессе, в т. ч. этические аспекты и риски, на уровне официальных документов.

В данной работе были поставлены следующие исследовательские вопросы:

- 1) Каково общее отношение российских вузов к использованию ИИ и как это отражено в их политиках?
- 2) Какие конкретные практические рекомендации по использованию ИИ даются ключевым участникам образовательного процесса – студентам и преподавателям?
- 3) Какие этические аспекты и потенциальные риски использования ИИ рассматриваются и регламентируются в официальных документах вузов?

Обзор литературы

В научной литературе представлен обширный пласт исследований, в которых отношение к интеграции ИИ в образование рассматривается сквозь призму двух противоположных позиций – пессимистической, акцентирующей внимание на рисках и вызо-

вах, и оптимистической, делающей акцент на перспективах и возможностях.

Сторонники первой, пессимистической, позиции сосредотачивают внимание на угрозах и рисках генеративных моделей ИИ в образовании. В категориях академической честности исследователи описывают плагиат, списывание, а также разрушение валидности оценивания, поскольку традиционные формы проверки знаний становятся уязвимыми перед возможностями ИИ [7; 8]. Категория когнитивного воздействия включает снижение критического мышления, когнитивную разгрузку [9]. Этические проблемы, в свою очередь, охватывают алгоритмическую предвзятость и предубежденность, заложенные в моделях, отсутствие прозрачности их работы [7; 9; 10]. В сфере конфиденциальности и безопасности ключевыми рисками выступают утечка данных и проблемы слежения [7]. Социально-экономические аспекты включают цифровое неравенство, углубляющее разрыв между разными группами студентов, замещение рабочих мест в образовательной сфере и монополизацию рынка образовательных технологий узким кругом крупных разработчиков. Психологические последствия выражаются в функциональной зависимости от ИИ, росте ИИ-тревожности, а также в сокращении полноценного человеческого взаимодействия в образовательном процессе. Наконец, экологическая категория фиксирует значительный экологический след, который оставляют вычислительные нагрузки, необходимые для обучения и функционирования ГИИ [9].

Сторонники второй позиции, оптимистической, анализируют потенциал генеративных моделей искусственного интеллекта в образовании, выделяя такие категории, как персонализация и адаптивное обучение, эффективность преподавания и административная поддержка, языковая поддержка и доступность, развитие когнитивных навыков и вовлеченность, демократизация доступа, а также поддержка психического благополучия [11–15]. Развитие когнитивных навыков

и вовлечённость обеспечиваются через интерактивность, геймификацию и мгновенную обратную связь, что усиливает мотивацию и развивает цифровую грамотность. Демократизация доступа снижает зависимость от традиционной инфраструктуры, предоставляя равные возможности студентам из удалённых регионов [11; 13–15].

Главным акцентом большинства современных исследований становится признание бесполезности запретительных мер в отношении использования ИИ в высшем образовании [16; 17]. Исследователи сходятся во мнении, что «гонка вооружений» между детекторами ИИ и языковыми моделями не имеет перспектив, а тотальные запреты не работают в условиях, когда более 90% студентов уже используют ИИ в своей учебной деятельности [18]. Вместо этого предлагается стратегия совместной выработки правил с участием политиков, исследователей и педагогов, направленная на создание прозрачных институциональных политик, охватывающих вопросы академической честности, оценки знаний и защиты данных [16; 17]. Массовое принятие ИИ студентами сочетается с неготовностью институциональных механизмов управления [18–20].

Особое значение в этой дискуссии приобретает вопрос о сохранении защитной функции образования. Традиционная организация обучения выполняет важную роль, противостоя хаосу «мгновенного обучения», который несет с собой ИИ. Регулирование в этом контексте понимается как процесс адаптации образовательной системы к новым технологическим реалиям при обязательном сохранении ее гуманистической основы [21].

Альтернативой ограничительному подходу выступает стратегия центрированной на студенте динамичной регламентации [22; 23]. Этот подход отвергает обе крайности, а именно хаотичное использование ИИ и тотальные запреты. Вместо этого он предлагает гибкую, постоянно адаптируемую политику, цель которой заключается не в контроле и наказании, а в обучении этичному и эффективному

использованию ИИ [22]. Ключевым элементом становится развитие алгоритмической грамотности, а институциональная политика должна учитывать, что присутствие ИИ ставит под сомнение традиционные методы оценивания: оценивается ли понимание студента или результат работы алгоритма [23].

Однако эмпирический анализ реальных политических документов университетов обнаруживает глубокую противоречивость институционального подхода. Даже в поощряющих политиках наблюдается внутренняя двойственность: одновременно предоставляются рекомендации по предотвращению использования ИИ студентами и признается ненадежность инструментов обнаружения. Регулирование часто перекладывает на преподавателей бремя пересмотра педагогических подходов без учета долгосрочных последствий, затрат времени и этических проблем. Это ставит под вопрос оптимистичный нарратив о «совместной выработке правил» и требует более пристального внимания к разрыву между декларируемыми принципами и практическими механизмами их реализации [24].

На этом фоне всё более значимым становится подход, смещающий фокус с внешних запретов и санкций на внутренние этические ориентиры. Согласно этой логике, личные этические принципы преподавателей играют более важную роль в ответственном использовании ИИ, чем формальные институциональные политики [25]. Устойчивое регулирование должно опираться на этические руководства и образование в области ответственного использования, а не на ограничения. Ключевым становится формирование у студентов способности к этическому суждению [26; 27], а не просто соблюдению правил.

Институциональная этика в этом контексте понимается не как абстрактная ценность, а как конкретное дизайн-решение. Когда институты полагаются на метрики и автоматизированные системы без этического сопровождения, они учат студентов «исполнять» добросовестность, а не практиковать ее. Регулирование должно включать создание условий

для формирования этического суждения, а не только систему санкций [28]. Необходим переход от абстрактных ценностных деклараций к конкретным сценариям использования, риск-ориентированным подходам и четкому распределению ответственности [29].

В российском академическом дискурсе исследуются в основном те же риски и угрозы, что и за рубежом [30–32]. Уникальной чертой российской дискуссии выступает акцент на угрозе унификации и формализации знаний, ведущей к потере плюрализма познавательных подходов, а также более острая постановка проблемы дегуманизации образования и утраты его воспитательной функции [33; 34]. В то же время категории, представленные в зарубежной литературе, – экологический след вычислительных нагрузок и монополизация рынка образовательных технологий – практически отсутствуют.

В отношении преимуществ российские и зарубежные подходы схожи [30; 35; 36]. Уникальной чертой российской дискуссии является акцент на роли ИИ в научно-исследовательской деятельности и подготовке кадров высшей квалификации, а также на цифровизации как инструменте сохранения единства образовательного пространства [33; 35].

При этом исследователи разделяются на сторонников эволюционного пути (постепенная смена парадигмы с компетентностного подхода на творчески ориентированный, формирование новых этических норм, изменение роли преподавателя) и сторонников революционного сценария (кардинальное сужение учебных программ, оставляющее только те курсы, где преподаватель обладает опытом, превосходящим возможности ИИ) [37; 38].

В сфере регулирования оба дискурса сходятся в признании бесполезности запретительных мер и необходимости пересмотра подходов к преподаванию и оценке знаний [31; 32; 36]. Российские авторы делают более выраженный акцент на формально-юридическом регулировании на уровне федерального законодательства [30; 34; 39]. В зару-

бежной литературе акцент смещён в иную плоскость. Вместо формализации права на обжалование алгоритмических решений на законодательном уровне исследователи сосредотачиваются на институциональных политиках университетов и стратегии «центрированной на студенте динамичной регламентации».

В российской литературе практически отсутствует рефлексия о разрыве между декларируемыми принципами и практическими механизмами реализации, который зарубежные исследователи фиксируют как ключевую проблему институциональных политик. Предложения по количественным ограничениям, остаются единичными и не формируют системного подхода [40]. Даже в вузах, где такие нормы вводятся, сохраняется устойчивый «нормативный разрыв» между повседневной практикой использования ИИ студентами и этическими оценками его приемлемости, что свидетельствует о том, что локальные количественные ограничения не восполняют отсутствие комплексной институциональной политики, а лишь подчеркивают фрагментарность подхода в российском контексте [41].

Методы

Исследование построено на всестороннем анализе документов, находящихся в открытом доступе на официальных сайтах университетов. В качестве основного исследовательского метода применяется качественный контент-анализ, позволяющий выявить и систематизировать темы и концепции в текстовых материалах. Как отмечают М. Стейси и Н. Моклер, анализ образовательной политики требует не просто описания содержания документов, но и понимания того, как эти документы конструируют реальность, какие дискурсы легитимизируют и какие ценности транслируют. В этой традиции контент-анализ нормативных документов рассматривается не как техническая процедура, а как теоретически обоснованный метод, позволяющий выявить как

явные, так и неявные смыслы политических текстов [42]. К. Крипендорф указывает, что основной задачей контент анализа является выявление скрытых закономерностей, намерений, ценностей и эффектов, которые эти документы порождают в обществе [43].

В фокус анализа попадают разнообразные форматы документов, непосредственно касающиеся регулирования искусственного интеллекта в академической среде: внутренние политики и стратегии, административные регламенты и процедуры, декларации о принципах, а также практические руководства и инструкции, связанные с использованием технологий ИИ в учебной и научной деятельности.

Для обеспечения репрезентативности и объективности выборки был применён системный подход к отбору учебных заведений. В основу лёг глобальный сводный рейтинг вузов России, который был выбран по причине комплексности методики его построения¹.

Для детального изучения были отобраны первые 97 ведущих вузов, которым в рейтинге присвоены высшие категории А+ и А. Эти университеты как правило, являются флагманами системы высшего образования, задающими институциональные тренды, и обладают наибольшими ресурсами для разработки и внедрения современных решений. С точки зрения неoinституциональной теории, именно эти вузы с наибольшей вероятностью выступают «ранними последователями» институциональных инноваций и образцами для региональных вузов.

Процедура сбора данных для максимального охвата была организована двумя основными способами: 1) с использованием поисковых систем (Яндекс, Google) по целевым запросам с указанием названия университета и ключевых терминов («искусственный интеллект», «генеративный искусственный интеллект», «кодекс этики ИИ», «этические принципы»; 2) путём прямой ручной навига-

ции по разделам официальных сайтов вузов: «Документы» или «Локальные нормативные акты», «Академическая политика», «Организация образовательного процесса», «Методические рекомендации», а также разделам, содержащим требования к оформлению и написанию выпускных квалификационных работ и диссертаций и разделам о системе «Антиплагиат. ВУЗ». Период сбора данных: декабрь 2025 – январь 2026 гг.

В ходе поиска только в 10 университетах из 97 были найдены документы, регламентирующие использование ИИ в образовании. Среди них представлены как столичные вузы (Национальный исследовательский институт Высшая школа экономики (НИУ ВШЭ), Российская академия народного хозяйства и государственной службы при Президенте РФ (РАНХиГС), Российский университет дружбы народов имени Патриса Лумумбы (РУДН), Московский государственный педагогический университет (МГПУ), Национальный исследовательский институт «МЭИ» (НИУ «МЭИ»)), так и ведущие региональные университеты (Томский государственный университет (ТГУ), Уральский федеральный университет имени первого Президента России Б.Н. Ельцина (УрФУ), Казанский национальный исследовательский технологический университет (КНИТУ)), а также классические университеты (Санкт-Петербургский государственный университет (СПбГУ)) и специализированный медицинский вуз (Рязанский государственный медицинский университет имени академика И.П. Павлова (РязГМУ)).

Для качественного контент-анализа политик были разработаны три аналитические категории, отражающие предложенную теоретическую рамку регулирования ИИ в высшем образовании: (1) степень проработанности правил (уровень детализации требований и процедур), (2) степень жесткости ограничений (доля запретительных

¹ Глобальный сводный рейтинг вузов России 2026 . URL: <https://tabiturient.ru/globalrating/> (дата обращения: 17.02.2026).

формулировок и санкций) и (3) ориентир на активную интеграцию (наличие стимулирующих мер и поддерживающих практик). Две первые категории операционализировались как порядковые шкалы с заранее заданными критериями от низкого к высокому уровню, третья – как бинарная (наличие/отсутствие признака) и применялись к полному корпусу документов до начала анализа результатов.

Кодирование текста политик осуществлял один исследователь. Такой дизайн соответствует распространённым протоколам качественного контент-анализа для *solo-researcher*, когда целью является концептуально насыщенное, а не статистически обобщаемое описание поля.

Критерий степени проработанности правил показывает, насколько подробно университет описал использование генеративного ИИ:

А) *Очень высокая степень*. Есть отдельный документ (или пакет документов), полностью посвященный ИИ. Описаны этические принципы, допустимые и запрещенные сценарии для всех категорий участников. Разграничено применение ИИ по видам работ и ролям, даны подробные рекомендации по использованию.

Б) *Высокая степень*. Есть отдельный документ по ИИ, но он односторонний: либо детально про академическую честность и техническое применение ГИИ без этики, либо подробно про этику без конкретных сценариев использования. При этом прописано разделение по видам работ и ролям и есть рекомендации по использованию.

В) *Средняя степень*. Единой университетской политики нет, использование ИИ регулируется фрагментарно. Правила общие и краткие, без детализации по видам работ и без четкого распределения ответственности.

Г) *Низкая степень*. Собственной политики по ИИ нет. ИИ упоминается лишь эпизодически в других документах (например, про «Антиплагиат» или ГИА).

Критерий степени жесткости ограничений показывает, насколько строг университет и каковы санкции:

А) *Высокая жесткость*. Четко прописаны санкции (вплоть до недопуска к защите, аннулирования результата). Подробно описаны механизмы выявления нарушений и процедура декларирования использования ИИ.

Б) *Средняя жесткость*. Запреты и санкции сочетаются с рекомендациями и разъяснениями. Последствия есть, но соразмерны проступку и допускают исправление. Акцент на обучении правильному использованию, за грубые/систематические нарушения предусмотрены дисциплинарные меры.

В) *Низкая жесткость*. Университет делает ставку на доверие и саморегуляцию. В документах акцент на этике и рекомендациях, санкции либо минимальны, либо не описаны. Механизмы контроля отсутствуют, ответственность в основном лежит на самих участниках образовательного процесса.

Наличие цели на активную интеграцию демонстрирует, что вуз стремится не только к регулированию использования ИИ, а намерен целенаправленно встроить технологии ИИ в образовательный процесс, например, меняя методики преподавания или повышая квалификацию участников образовательного процесса в области использования ИИ.

Результаты исследования

Первый этап контент-анализа

Только в НИУ ВШЭ (включая филиалы) разработаны отдельные комплексные документы: «Декларация этических принципов»², «Правила использования ИИ студентами»³

² Декларация этических принципов создания и использования систем искусственного интеллекта в Национальном исследовательском университете «Высшая школа экономики» // Сайт НИУ ВШЭ. – URL: <https://www.hse.ru/docs/969670638.html> (дата обращения: 17.02.2026).

³ Правила использования искусственного интеллекта студентами НИУ ВШЭ // Сайт НИУ ВШЭ. – URL: https://www.hse.ru/studyspravka/ai_guidelines (дата обращения: 17.02.2026).

и «Регламент проверки работ»⁴. Задана этическая рамка использования ИИ, выведены правила использования студентами ГИИ в учебном процессе, а также разработан регламент проверки письменных студенческих работ на использование ГИИ студентами.

В РАНХиГС⁵, СПбГУ, РУДН, ТГУ и УрФУ имеются утвержденные приказами политики или комплексные документы, регулирующие использование на разных уровнях образовательной деятельности. Особенно важно выделить ТГУ и УрФУ, в документах которых четко определены и описаны этические принципы ИИ в образовании.

В КНИТУ и НИУ «МЭИ» утверждены положения, регламентирующие использование системы «Антиплагиат. ВУЗ», которые содержат прямые указания на проверку работ на наличие текста, сгенерированного ИИ, и определяют последствия его незадекларированного использования. В МГПУ использование средств ГИИ при написании ВКР регламентировано пунктом в Положении о проведении государственной итоговой аттестации.

В РязГМУ действует положение об учебных изданиях, устанавливающее допустимый процент использования алгоритмов ИИ при их подготовке.

Анализ документов позволяет выявить два основных подхода к определению круга лиц, на которых распространяются данные правила.

1. Три университета (НИУ ВШЭ, ТГУ, УрФУ) формулируют политики, охватывающие все категории участников образовательного процесса. Правила и принципы применения ИИ адресованы одновременно

профессорско-преподавательскому составу, администрации, студентам, а иногда и третьим лицам. Целью данных документов является установление единых рамок для всей университетской экосистемы.

2. Другая группа вузов (7 университетов) приняла документы, которые в первую очередь направлены на регулирование деятельности студентов. В центре внимания находится использование ГИИ. Основная задача таких документов – дать четкие инструкции и ограничения для обучающихся при выполнении работ с помощью ИИ.

Общее отношение вузов к ГИИ: поощрение, ограничение, запрет

Анализ документов демонстрирует, что абсолютный запрет на использование ГИИ не зафиксирован ни в одном из 10 проанализированных документов. ГИИ признаётся неотъемлемым элементом образовательной среды, и его интеграция в академические процессы происходит путём разработки соответствующей нормативной базы.

Ни в одном из рассмотренных документов не выявлено и абсолютного поощрения использования ГИИ студентами и преподавателями. Доминирует подход, сочетающий активную интеграцию и регуляцию. Однако вузы демонстрируют существенные различия как в степени проработанности правил, так и в степени жесткости ограничений (Табл. 1).

Критерий степени проработанности правил не требует отдельного детального рассмотрения, однако два других критерия следует раскрыть через конкретные примеры более подробно.

Высокая степень жесткости ограничений выявлена лишь в одном университете

⁴ Регламент организации проверки письменных учебных работ на наличие плагиата, использования генеративных моделей и размещения выпускных квалификационных работ обучающихся по программам бакалавриата, специалитета и магистратуры на корпоративном сайте НИУ «Высшая школа экономики» // Сайт НИУ ВШЭ. – URL: <https://www.hse.ru/docs/922831988.html> (дата обращения: 17.02.2026).

⁵ Политика академической честности и использования искусственного интеллекта в ФГБОУ ВО «Российская академия народного хозяйства и государственной службы при Президенте Российской Федерации» // Сайт РАНХиГС. – URL: https://www.ranepa.ru/upload/doc/prikazy-ranhigs/2025/02-1062_10.06.2025.pdf (дата обращения: 17.02.2026).

Таблица 1

Степень проработанности правил и жесткости ограничений использования ГИИ в учебном процессе

Table 1

Level of Regulatory Detail and Restrictiveness regarding GenAI in Education

Университет	Степень проработанности правил	Степень жесткости ограничений	Ориентир на активную интеграцию
НИУ ВШЭ	Очень высокая	Высокая	да
УрФУ	Очень высокая	Средняя	да
ТГУ	Высокая	Низкая	да
СПбГУ	Средняя	Низкая	да
РАНХиГС	Высокая	Средняя	да
РУДН	Высокая	Средняя	нет
МГПУ	Низкая	Низкая	нет
НИУ «МЭИ»	Низкая	Средняя	нет
КНИТУ	Низкая	Средняя	нет
РязГМУ	Низкая	Низкая	нет

(НИУ ВШЭ). В регламенте, во-первых, прописывается процедура проверки как работ общего образовательного процесса, так и выпускной квалификационной работы (ВКР). Во-вторых, в документе не прописывается возможность процедуры исправления ситуации с использованием незадекларированного ИИ, а сразу указываются санкции в следующих формулировках: «В случае выявления плагиата или незадекларированного использования генеративных моделей в письменной учебной работе студенту (студентам) выставляется оценка «0» за соответствующий элемент контроля. Кроме того, в случае выявления плагиата в письменной учебной работе руководителем работы может быть инициирована процедура привлечения студента (студентов) к дисциплинарной ответственности за нарушение академических норм».

К средней степени жесткости ограничений отнесены пять университетов. И, хотя санкции присутствуют, университеты оставляют пространство для исправления ситуации и стремятся научить правильному использованию ИИ-инструментов. В УрФУ⁶ ограничения на использование ГИИ прописаны как для студентов, так и для преподавателей и администрации. В качестве санкции для всех участников указано: «Нарушение политики влечёт предусмотренную законодательством РФ ответственность и дисциплинарные меры», без каких-либо подробностей. В РУДН⁷ акцентируют внимание на обязанности студента указывать, если ГИИ применялся в теоретической или эмпирической частях работы. В РАНХиГС за нарушения предусмотрены устные или письменные предупреждения, снижение академической оценки (аннулирование

⁶ Политика в области внедрения и использования систем искусственного интеллекта в образовательной, научной и управленческой деятельности УрФУ // Сайт УрФУ. – URL: https://urfu.ru/fileadmin/user_upload/urfu.ru/documents/education/Prikaz_po_osnovnoi_deyatelnosti_No0896_03_ot_02.10.2025_O_vvedenii_v_deistvie_Politiki_v_oblasti_v3_.pdf (дата обращения: 17.02.2026).

⁷ Политика РУДН в области применения обучающимися генеративного искусственного интеллекта при подготовке (написании/выполнении) учебных (учебно-научных) работ // Сайт РУДН. – URL: <https://www.rudn.ru/storage/media/documents/7c62d88a-8d8e-4e3a-a2a0-9fcc74c3e12d/uxxGZjVlg0ESgNkz7EuwgpQ9D2ToEOPT3V9QL1O5.pdf> (дата обращения: 17.02.2026).

баллов), дисциплинарные взыскания (замечание, выговор, а также отчисление, применяемое только в крайних случаях за наиболее грубые нарушения), и проводятся воспитательные меры, направленные на формирование нетерпимого отношения к недобросовестности.

В НИУ «МЭИ»⁸ и КНИТУ⁹ документы сосредоточены только на подготовке к ВКР и выявлении использования студентами ГИИ системой «Антиплагиат». В этих вузах подход практически не различается. При выявлении факта использования ИИ для написания ВКР он фиксируется в отчёте и студенту предоставляется возможность устранения плагиата, и, только если повторная проверка обнаружит факт использования ГИИ, студент подлежит отчислению.

Четыре университета были отнесены к образовательным организациям с низкой степенью жесткости ограничений. В ТГУ¹⁰ акцент делается на этических принципах и санкций за их нарушение не прописано. В СПбГУ¹¹ лишь предлагается ограничивать объем сгенерированного контента – до 20% от всей работы обучающегося, однако о последствиях нарушения рекомендации

в документах не сказано. В МГПУ¹² указывают на обязанность декларирования использования ГИИ, но также не говорят о последствиях нарушения. В РязГМУ¹³ также никакой информации о санкциях нет. В этой категории вузов ответственность за соблюдение правил в большей степени лежит на самом студенте, преподавателе или исследователе.

Наличие цели на активную интеграцию демонстрирует, что университет намерен целенаправленно встроить технологии ИИ в образовательный процесс. К примеру, в РАНХиГС прописывают, что педагогическим работникам рекомендуется: «Включать ИИ в обучение. Давать задания, предполагающие использование ИИ, но требующие вовлечения обучающихся в работу. Включать в обучение командную работу и групповые проекты обучающихся с применением ИИ». В ТГУ прописано, что «политика направлена на просвещение, обучение и популяризацию знаний о технологиях ИИ, а также формирование должных навыков работы с инструментами ИИ у всех участников образовательного процесса». УрФУ указывает, что политика в области

⁸ Положение об использовании системы «Антиплагиат. ВУЗ» для оценки на заимствования письменных учебных и научных работ обучающихся по программам бакалавриата, программам специалитета и программам магистратуры в ФГБОУ ВО НИУ «МЭИ» // Сайт НИУ МЭИ. – URL: <https://mpei.ru/AboutUniverse/OfficialInfo/Orders2025/MPEI-25-587.pdf> (дата обращения: 17.02.2026).

⁹ Положение о порядке использования системы «Антиплагиат.ВУЗ»: утв. приказом ФГБОУ ВО «КНИТУ» от 02.04.2025 № 334-0 // Сайт КНИТУ. – URL: <https://www.kstu.ru/servlet/contentblob?id=465197> (дата обращения: 17.02.2026).

¹⁰ Политика использования искусственного интеллекта в образовательном процессе ФГАОУ ВО «Национальный исследовательский Томский государственный университет» // Сайт ТГУ. – URL: <https://clcl.i/fjDUp> (дата обращения: 17.02.2026).

¹¹ Политика применения инструментов генеративного искусственного интеллекта (ИИ) в процессе обучения, подготовки курсовой и выпускной квалификационной работы учащимися, исследователями и преподавателями ВШМ СПбГУ // Сайт СПбГУ. – URL: <https://gsom.spbu.ru/about-gsom/dokumenty/politic/> (дата обращения: 17.02.2026).

¹² О внесении изменений в Положение о проведении ГИА по образовательным программам высшего образования – программам бакалавриата, программам специалитета, программам магистратуры: приказ ГАОУ ВО МГПУ от 4.09.2023 № 633 // Сайт МГПУ. – URL: https://www.mgpu.ru/wp-content/uploads/2023/09/04.09.2023_633obshh_O-vnesenii-izmenenij-v-GIA.pdf (дата обращения: 17.02.2026).

¹³ Положение об учебных и методических изданиях: утв. приказом ФГБОУ ВО РязГМУ Минздрава России от 01.10.2025 № 577-д // Сайт РязГМУ. – URL: https://www.rzgm.ru/sveden/files/viq/Pologhenie_ob_uchebnyx_i_metod_izd.pdf (дата обращения: 17.02.2026).

ИИ направлена на «создание условий для широкого внедрения в деятельность УрФУ прорывных технологий на основе систем искусственного интеллекта».

С позиций неинституциональной теории, отсутствие запретов объясняется действием логики уместности. В современном глобальном академическом дискурсе полный запрет ИИ воспринимался бы как нелегитимный и противоречащий имиджу передового университета.

Рекомендации по использованию генеративного ИИ для студентов и преподавателей
Для студентов:

1. Обязательное декларирование. Требование указывать факт, цели, инструменты и объем использования ИИ в работе (в тексте, сносках или специальной форме, как в НИУ ВШЭ).

2. Ограничение объема и цели использования. Запрет на полное написание ключевых разделов (введение, выводы, расчеты – НИУ «МЭИ», РУДН). Рекомендация ограничивать долю ИИ-контента (СПбГУ – до 20%). Запрет на использование для прохождения итоговой аттестации (УрФУ).

3. Критическая проверка. Обязательная верификация всех фактов, данных и источников, полученных от ИИ (УрФУ, МГПУ, РАНХиГС, СПбГУ).

4. Рекомендации по использованию ИИ для вспомогательных задач: поиск информации, анализ данных, генерация идей, редактирование текста, составление плана (УрФУ, МГПУ, РУДН). МГПУ предоставляет конкретные примеры промптов.

5. Соблюдение конфиденциальности. Запрет на загрузку в публичные ИИ-системы персональных данных, конфиденциальной информации (УрФУ, ТГУ, РАНХиГС).

Для преподавателей и научных руководителей:

1. Разработка локальных правил. СПбГУ рекомендует преподавателям создавать па-

мятки или включать правила по ИИ в рабочие программы дисциплин.

2. Контроль и оценка. Преподаватель имеет право запретить использование ИИ для конкретных заданий. На преподавателей возлагается обязанность проверять работы на наличие незадекларированного ИИ-контента и принимать решение о правомерности его использования (НИУ ВШЭ, НИУ «МЭИ»).

3. Методическая и воспитательная работа. РАНХиГС и РУДН подчеркивают необходимость обучения студентов этическому использованию ИИ.

4. Норматив для авторов учебных изданий. РязГМУ устанавливает для преподавателей количественный лимит использования алгоритмов ИИ при подготовке учебных и учебно-методических изданий – не более 5%.

Содержание этических принципов в документах

Этические принципы использования ИИ сформулированы и прописаны только в документах трёх университетов – НИУ ВШЭ, ТГУ, УрФУ. Это указывает на то, что детальная этическая рефлексия пока не стала обязательным элементом легитимации для большинства вузов. Наличие чётко сформулированных принципов – признак более продвинутой стадии институционализации, когда организация не просто реагирует на проблему, а выстраивает целостную нормативную рамку, интегрированную в глобальный этический дискурс.

При этом такие университеты, как РАНХиГС и СПбГУ, хотя и не определяют прямо этические принципы, при разработке своих политик опирались на общероссийские документы. Ключевым ориентиром для них служит Кодекс этики в сфере искусственного интеллекта, разработанный Альянсом в сфере искусственного интеллекта¹⁴. Кроме него, учитывались Нацио-

¹⁴ Кодекс этики в сфере искусственного интеллекта // Сайт Альянса в сфере искусственного интеллекта. – URL: <https://ethics.a-ai.ru/> (дата обращения: 17.02.2026).

нальная стратегия развития искусственного интеллекта до 2030 года¹⁵, Руководство по использованию генеративного искусственного интеллекта в образовании и научных исследованиях ЮНЕСКО¹⁶ и другие материалы¹⁷. Несмотря на то, что университеты начали регулировать использование ИИ и ориентируются на общенациональные этические стандарты, в большинстве случаев эти принципы не прописаны в их локальных документах прямо и развернуто. Фактически, прямое формулирование этики ИИ остается исключением, а не правилом. Это можно интерпретировать как церемониальное принятие норм. Университет демонстрирует, что этические нормы необходимы, но не инвестирует ресурсы в детальную проработку, возможно, потому что это пока не критично для легитимности или потому что реальные практики ещё не устоялись.

Второй этап контент-анализа

На этом этапе выборка была расширена за счет документов еще трех университетов (Пензенский государственный университет (ПГУ), Ростовский государственный экономический университет (РГЭУ (РИНХ)), Южно-Уральский государственный университет (национальный исследовательский университет) (НИУ ЮУрГУ)). Они были найдены по поисковым запросам: «Кодекс этики в вузах РФ» и «Политика использования ИИ в образовании» в браузерах Яндекс и Google.

Сравнительный анализ формулировок этических принципов выявил расхождения в смысловых акцентах при сохранении общей концептуальной основы. Данные различия отражают процесс адаптации глобальных этических дискурсов к специфике локальных контекстов: каждый университет адап-

тирует общие принципы к своим организационным потребностям, ресурсам и представлениям о рисках.

Все рассматриваемые вузы выделяют несколько основных групп принципов и рисков, что позволяет говорить об общей согласованности, несмотря на некоторые различия в их интерпретации.

Приоритет человека и человеческого общения

Все университеты признают риск вытеснения общения и подмены роли преподавателя, позиционируя ИИ как вспомогательный инструмент для развития человека. Этот принцип защищает профессиональную идентичность преподавателя и легитимизирует традиционную педагогическую модель в условиях проникновения ИИ. Так, в политике ЮУрГУ зафиксирован «приоритет человеческого общения», ПГУ рассматривает ИИ «исключительно как вспомогательный инструмент», а НИУ ВШЭ подчеркивает, что ИИ «должен способствовать расширению возможностей межличностного взаимодействия».

Академическая честность и прозрачность

Всеми университетами признаётся риск нарушения принципов академической честности и прозрачности. Поэтому учебные заведения в рамках этических принципов устанавливают требования о декларировании и маркировке использования ИИ в учебных работах студентов. А с точки зрения прозрачности фиксируется необходимость равной информированности о возможностях, ограничениях и рисках использования ИИ. Эти принципы являются механизмом институционального контроля, который делает использование ИИ ви-

¹⁵ О развитии искусственного интеллекта в Российской Федерации. Указ Президента РФ № 490 (ред. от 15.02.2024) // Гарант.Ру. – URL: https://base.garant.ru/72838946/#block_1000 (дата обращения: 17.02.2026).

¹⁶ ЮНЕСКО. Руководство по использованию генеративного искусственного интеллекта в сфере образования и научных исследований. Париж, 2023. 44 с. // Сайт ЮНЕСКО. – URL: <https://unesdoc.unesco.org/ark:/48223/pf0000386693> (дата обращения: 17.02.2026).

¹⁷ Проект Кодекса этики в сфере ИИ в образовании // Сайт Альянса в сфере искусственного интеллекта. – URL: <https://ethics.a-ai.ru/ethics-of-education/> (дата обращения: 17.02.2026).

димым и подотчётным. Например, в РГЭУ (РИНХ)¹⁸ требование сформулировано предельно четко: «Любое использование ИИ в работе должно быть явно и полностью декларировано. Участник обязан указывать, какие именно части были созданы или существенно изменены с помощью ИИ». В НИУ ВШЭ также указывают, что работники и обучающиеся «обязаны выделять результаты своей деятельности, при реализации которой был использован ИИ, указывая характер и объем работ, выполненных с его помощью».

Безопасность и конфиденциальность

Университетами отмечаются риски, которые связаны с безопасностью, обработкой персональных данных, нарушением конфиденциальности и авторских прав. В принципах указывается необходимость ответственного использования систем ИИ. Эти принципы также встраивают регулирование ИИ в более широкий контекст защиты данных и соблюдения прав человека, что повышает легитимность политик в отношении ИИ. В ЮУрГУ¹⁹ отмечается: «Внедрение и применение ИИ... должны быть безопасными и обеспечивать защиту персональных данных». Сходную позицию занимает ПГУ²⁰: «При использовании технологий ИИ запрещается вводить в составе запроса персональные данные, конфиденциальную информацию и информацию, нарушающую авторские права». В НИУ ВШЭ добавляют важный нюанс: «Не следует вводить персональные данные, конфиденциальную информацию, интеллектуальную собственность третьих лиц на

какую-либо платформу ИИ... если это приводит к нарушению законодательства».

Равный доступ

Указывается и риск усиления социального неравенства и дискриминации, во-первых, из-за предвзятости самих алгоритмов и данных, а во-вторых, из-за финансовых барьеров и возможностей получения доступа к платным сервисам. Поэтому вузы указывают, что соблюдение данного принципа «справедливости и недискриминации» найдется в рамках их ответственности. В формулировках принципов указывается, что для недопущения неравенства предусмотрены консультации, поддержка по применению инструментов, а также обучение эффективному применению ИИ. В РГЭУ (РИНХ) подчеркивают: «Условия для использования ИИ должны быть одинаковыми для всех участников... независимо от их финансового положения или уровня владения технологиями». В ЮУрГУ обращают внимание на двойственный характер проблемы: «Применение ИИ не может быть выдвинуто в качестве обязательного требования... если влечет за собой личные финансовые траты. Также необходимо учитывать возможность возникновения ограничений... вследствие которых участники... могут оказаться в неравном положении». В НИУ ВШЭ добавляют важный организационный механизм: «В случае необходимости его [ИИ] использования администрация... принимает меры к тому, чтобы преподаватели и обучающиеся имели к нему равный доступ». Что касается дискриминации, то в документах ТГУ указывается:

¹⁸ Политика использования технологий генеративного искусственного интеллекта в образовательной и научно-исследовательской деятельности в ФГБОУ ВО «РГЭУ (РИНХ)» // Сайт РГЭУ (РИНХ). – URL: https://rsue.ru/universitet/dokumenty/doc/Politika_iskust_intilekta.pdf (дата обращения: 17.02.2026).

¹⁹ Политика в области применения искусственного интеллекта в процессе обучения и научно-исследовательской работе студентами и педагогическими работниками в ФГАОУ ВО «Южно-Уральский государственный университет (национальный исследовательский университет)» // Сайт ЮУрГУ. – URL: <https://clcl.i/BZNRh> (дата обращения: 17.02.2026).

²⁰ Политика в области применения искусственного интеллекта в процессе обучения и научно-исследовательской работе студентами и педагогическими работниками в ФГАОУ ВО «Южно-Уральский государственный университет (национальный исследовательский университет)» // Сайт ЮУрГУ. – URL: <https://clcl.i/BZNRh> (дата обращения: 17.02.2026).

«необходимо проводить оценку результатов, полученных от систем ИИ, через призму общечеловеческих моральных принципов и ценностей», а в НИУ ВШЭ поясняют, что «ИИ обучается на основе данных, созданных людьми, которые могут тиражировать предвзятость и социальные стереотипы, способствующие дискриминации», поэтому необходимо «предпринимать усилия для их предотвращения».

Наряду с общими этическими принципами были выделены уникальные – принцип стохастичности и принцип разумного ограничения, что явно указывает на разные подходы в стратегиях регулирования ИИ. Стоит отметить и то, что уникальные принципы демонстрируют способность вуза выходить из режима институционального изоморфизма и становиться генератором новых норм.

Их анализ позволил выделить два подхода в стратегиях регулирования:

1. Адаптивный подход (основан на принципе стохастичности).

Технологии ИИ меняются быстрее, чем успевает сформироваться нормативная база, поэтому жёсткие и долгосрочные запреты неэффективны. В рамках данного подхода регулирование ИИ заключается в установлении гибких принципов, которые задают общее направление, но оставляют пространство для внесения изменений или пересмотра правил. Ключевыми механизмами становятся непрерывный мониторинг развития технологий, стимулирование внутренних исследований и открытых дискуссий в академической среде о последствиях их применения.

Этот подход отражает признание неопределённости и отказ от полного контроля. Он характерен для организаций, которые стремятся сохранить гибкость в условиях быстрых изменений. Принцип стохастичности легитимизирует незавершённость и временный характер правил, защищая университет от критики за несовершенство регулирования.

2. Управленческо-ограничительный подход (основан на принципе разумного ограничения).

Этот подход сфокусирован на оперативном реагировании на риски, связанные с ИИ. Принцип легитимизирует право администрации накладывать целевые запреты. В отличие от предыдущего подхода, здесь делается акцент на контроле. Подход носит более директивный характер и ориентирован на поддержание академических стандартов и порядка через администрирование.

Этот подход отражает логику рационального контроля и иерархического управления. Он характерен для организаций, которые выделяют основным приоритетом минимизацию рисков и защиту репутации над гибкостью. Это способ сохранить контроль над образовательным процессом в условиях технологической неопределённости.

Обсуждение

Выявленная нами картина низкого уровня институционализации использования ИИ согласуется с данными анализа 15 документов из восьми европейских стран, который показал, что регулирование ИИ в высшем образовании Европы находится на начальной стадии, а на уровне университетов сохраняется «лоскутная» картина, когда лишь немногие вузы разработали и опубликовали собственные руководства [44]. Другое исследование демонстрирует институциональный разрыв университетов в Центральной Азии. Из 30 рассмотренных учеными университетов лишь 2 (6,6%) имели опубликованные политики или руководства по использованию генеративного ИИ [45]. Аналогичные данные показывает исследование среди 500 ведущих университетов мира: менее трети (26,5%) имели формальные политики в отношении использования генеративного ИИ. Данные работы подтверждают, что низкий уровень институционализации является глобальной закономерностью, а не локальной особенностью [46]. Это согласуется и с выводами российских исследователей, которые также констатировали отсутствие единого стандарта и «базовый» характер существующих документов в ограниченном круге российских

вузов (МГПУ, НИУ ВШЭ, ТГУ) [40]. Наше исследование масштабирует этот вывод: даже наличие разрозненных документов является скорее исключением, чем тенденцией.

Однако, если в странах Западной Европы и США низкий уровень формальной регламентации проявляется в виде фрагментарности и «лоскутности» политик при высокой инициативе снизу (со стороны факультетов, преподавателей и администраций) [17; 47], а в Японии и Китае – в разрыве между государственными стратегиями и университетскими практиками [17; 48], то в России отсутствие формальных политик не компенсируется ни устойчивой низовой активностью, ни последовательным государственным регулированием, что делает институциональный вакуум более глубоким.

Настоящее исследование демонстрирует, что ни один из проанализированных российских вузов не ввел абсолютного запрета на использование генеративного ИИ. Доминирует модель «разрешено при условиях». Этот вывод согласуется с общемировой тенденцией, зафиксированной в исследованиях, где отмечается, что большинство университетов отказываются от тотальных запретов, признавая их неэффективность в условиях массового использования ИИ студентами [17; 49; 50].

Регулирование академической честности при использовании генеративного ИИ составляет ядро всех проанализированных политик: в них обычно закреплены требования к декларированию, ограничения на использование ИИ в ключевых разделах работ и механизмы проверки. Это согласуется с существующими исследованиями, которые показывают, что университетские политики в первую очередь сосредоточены на честности, этичном использовании ИИ и правилах цитирования [24; 51].

Исследование Института образования НИУ ВШЭ зафиксировало шесть моделей поведения вузов по отношению к генера-

тивному ИИ – от активного внедрения до полного игнорирования. Наши данные подтверждают эту вариативность: мы выделили три степени проработанности правил и три степени жесткости ограничений, что отражает дифференциацию институциональных стратегий. Как отмечает директор института Е. Терентьев, «когда университет не заявляет позицию по ИИ – это уже решение. И оно опасно»²¹. В нашем исследовании мы обнаружили, что почти 90% ведущих вузов России не имеют формальных политик, что можно интерпретировать как сознательное или вынужденное «решение не решать».

Выделенные нами степени проработанности правил могут служить «дорожной картой» для вузов, находящихся на разных этапах разработки собственных политик.

Однако данное исследование имеет ряд ограничений. Во-первых, оно ограничено временными рамками, представляя собой статичный срез, который не отражает динамику меняющейся сферы регулирования ИИ. Во-вторых, выборка исследования ограничена 10 вузами, что в совокупности с фокусом на лидерах рейтинга не позволяет распространить выводы на всю систему высшего образования РФ, при этом на втором этапе анализа выборка была расширена за счет вузов вне первоначального рейтинга, что нарушает изначальный принцип отбора и вносит элемент субъективности. В-третьих, метод сбора данных ограничен анализом только открытых источников, что могло привести к пропуску непубличных внутренних регламентов, а также к неполноте охвата документов в силу особенностей навигации на сайтах вузов. И, в-четвертых, сам качественный анализ определений этических принципов неизбежно содержит долю авторского восприятия.

Заключение

Обращение к типологии институционального изоморфизма позволяет объяснить

²¹ Вузы разделились на шесть лагерей в отношении к искусственному интеллекту // Сайт НИУ ВШЭ. – URL: <https://www.hse.ru/news/expertise/1074002356.html> (дата обращения: 17.02.2026).

причины создания университетами политик в отношении ИИ и одновременно разнородность этих политик. Миметический изоморфизм обнаруживается, вузы ориентируются на ограниченный круг лидеров и на внешние авторитетные источники, такие, например, как Кодекс этики Альянса в сфере искусственного интеллекта. Однако необходимо отметить, что прямых контекстуальных свидетельств сознательного воспроизведения формулировок из документов вузов-лидеров не было найдено. Это ограничение указывает на то, что миметический изоморфизм в российской действительности действует скорее на уровне заимствования жанра документа, чем на уровне осознанного копирования конкретного содержания, то есть перенимается общая структура и риторика.

Можно также сделать вывод о том, что нормативный изоморфизм проявился слабее всего. Об этом свидетельствует тот факт, что развёрнутые этические принципы прописаны лишь в трёх университетах первоначальной выборки.

Принудительный изоморфизм несколько парадоксален в российских реалиях. С одной стороны, существует национальная стратегия развития ИИ до 2030 года, а с другой – системные требования к вузам на федеральном уровне и даже региональных уровнях законодательства отсутствуют.

Обращение к модели трёх стадий институционализации позволяет оценить, на каком этапе находится процесс закрепления норм. Исследование показало, что экстернализация уже произошла – студенты массово используют генеративные модели искусственного интеллекта. Объективизация находится на ранней, фрагментарной стадии. Даже в тех десяти университетах, которые имеют документы в отношении регулирования ИИ, наблюдается значительная вариативность – от комплексных многоуровневых систем до единичных пунктов в положениях об антиплагиате или государственной итоговой аттестации, а этические принципы, как наиболее сложная форма объективизации,

присутствуют лишь в трёх случаях. Интернализацию, как усвоение установленных правил академическим сообществом и превращение их во внутренние нормы, в рамках данного исследования измерить не представляется возможным. Университеты реагируют на уже существующие практики, но не завершают цикл, что создает институциональный разрыв, при котором формальные правила существуют, но не интернализованы, а значит, не работают как «объективная реальность».

Говоря о перспективах исследования, можно выделить три направления будущих исследований. Первое – изучение процессов интернализации, а именно, как студенты и преподаватели в вузах с формальными политиками и без них воспринимают существующие правила, формируется ли у них внутреннее этическое суждение. Второе – сравнительный анализ с университетами из более низких рейтинговых категорий (В и С), чтобы проверить гипотезу о том, копируют ли они практики лидеров или остаются в состоянии институциональной инерции. Третье – лонгитюдное наблюдение за тем, как будет развиваться объективизация во времени: станут ли этические принципы обязательным элементом для всех университетов или же сохранится текущая дифференциация.

Литература

1. DiMaggio P.J., Powell W.W. The Iron Cage Revisited: Institutional Isomorphism and Collective Rationality in Organizational Fields // *American Sociological Review*. – 1983. – Vol. 48, no. 2. – P. 147–160. – URL: <http://www.jstor.org/stable/2095101> (дата обращения: 12.01.2026).
2. Pratt J. Institutional isomorphism and online learning in Australian higher education // *Inaugural Academy of World Business, Marketing and Management Development Conference*. – Academy of World Business, Marketing and Management Development, 2004. – URL: <https://opus.lib.uts.edu.au/bitstream/10453/7529/1/2004000718.pdf> (дата обращения: 12.01.2026).
3. Cardona Mejia L.M., Pardo del Val M., Dasi Coscollar A. The institutional isomorphism in

- the context of organizational changes in higher education institutions // *International Journal of Research in Education and Science (IJRES)*. – 2020. – Vol. 6, no. 1. – P. 61–73. – DOI: 10.46328/ijres.v6i1.639.
4. Baharuddin N.B., Abd Rahman N.H., Kamil N.L.M. Digital Education in Malaysian Public Universities: A Neo-Institutional Perspective on Policy Implementation // *International Journal of Modern Education*. – 2025. – Vol. 7, no. 27. – P. 571–588. – DOI: 10.35631/IJMOE.7270355.
 5. Немировский В.Г. Искусственный интеллект в зеркале неонинституционального подхода: социологический анализ // *Социальные и гуманитарные науки. Отечественная и зарубежная литература. Сер. 11, Социология*. – 2024. № 4. – DOI: 10.31249/rsoc/2024.04.03.
 6. Бергер П.А., Лукман Т. Социальное конструирование реальности: Договор о социологии знания : пер. с англ. / предисл. П. Бурдьё ; [пер. с англ. А.А. Назаренко]. – Москва: Академический проект, 2020. – 323 с. – ISBN: 5-85691-036-2.
 7. Owidi S.O., Lyanda J.N., Okono E.O., Micheni E.M. The Dark Side of Generative AI in Higher Education: Challenges, Ethical Concerns, and Implications // *East African Journal of Arts and Social Sciences*. – 2025. – Vol. 8, no. 3. – P. 526–543. – DOI: 10.37284/eajass.8.3.3690.
 8. Sharma R.C., Panja S.K. Addressing Academic Dishonesty in Higher Education: A Systematic Review of Generative AI's Impact // *Open Praxis*. – 2025. – Vol. 17, no. 2. – P. 251–269. – DOI: 10.55982/openpraxis.17.2.820.
 9. AlBlooshi S. Artificial intelligence in higher education, opportunities, and challenges: a review // *Frontiers in Education*. – 2026. – Vol. 10. – Article no. 1683968. – DOI: 10.3389/feduc.2025.1683968.
 10. Valentini A., Blancas A. The challenges of AI in higher education and institutional responses: Is there room for competency frameworks? UNESCO IESALC Working Paper. 2025. – URL: <https://www.iesalc.unesco.org/en/articles/challenges-ai-higher-education-and-imperative-competency-frameworks> (дата обращения: 17.02.2026).
 11. Jukiewicz M. How generative artificial intelligence transforms teaching and influences student wellbeing in future education // *Frontiers in Education*. – 2025. – Vol. 10. – Article no. 1594572. – DOI: 10.3389/feduc.2025.1594572.
 12. Luo J., Zheng C., Yin J., Teo H.H. Design and assessment of AI-based learning tools in higher education: a systematic review // *International Journal of Educational Technology in Higher Education*. – 2025. – Vol. 22. – Article no. 42. – DOI: 10.1186/s41239-025-00520-8.
 13. Zawacki-Richter O., Bai J.Y.H., Lee K., Slagter van Tryon P.J., Prinsloo P. New advances in artificial intelligence applications in higher education? // *International Journal of Educational Technology in Higher Education*. – 2024. – Vol. 21. – Article no. 32. – DOI: 10.1186/s41239-024-00471-6.
 14. Bond M., Khosravi H., De Laat M., Bergdahl N., Wegeia V. et al. A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigour // *International Journal of Educational Technology in Higher Education*. – 2024. – Vol. 21. – P. 1–24. – DOI: 10.1186/s41239-024-00467-2.
 15. Kalnina D., Nimante D., Baranova S. Artificial intelligence for higher education: benefits and challenges for pre-service teachers // *Frontiers in Education*. – 2024. – Vol. 9. – Article no. 1501819. – DOI: 10.3389/feduc.2024.1501819.
 16. Baidoo-Anu D., Owusu Ansah L. Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learning // *Journal of AI*. – 2023. – Vol. 7, no. 1. – P. 52–62. – DOI: 10.61969/jai.1337500.
 17. Li M., Xie Q., Enkhtur A., Meng S., Chen L. et al. A framework for developing university policies on generative AI governance: A cross-national comparative study // *arXiv*. – 2025. – DOI: 10.48550/arXiv.2504.02636.
 18. Isaifan R.J. Artificial intelligence in higher education: A global statistical synthesis for policy and quality assurance reform // *Education Sciences*. – 2026. – Vol. 16, no. 3. – Article no. 483. – DOI: 10.3390/educsci16030483.
 19. Karkoulou S., Sayegh N., Sayegh N. ChatGPT unveiled: Understanding perceptions of academic integrity in higher education—A qualitative approach // *Journal of Academic Ethics*. – 2025. – Vol. 23. – P. 1171–1188. – DOI: 10.1007/s10805-024-09543-6.
 20. Bernstein S., Rahman A., Sharifi N., Terbish A., MacNeil S. Beyond the benefits: A systematic review of the harms and consequences of generative AI in computing education // *arXiv*. – 2025. – URL: <https://arxiv.org/abs/2510.04443> (дата обращения: 17.02.2026).

21. Wu Y. Integrating generative AI in education: How ChatGPT brings challenges for future learning and teaching // *Journal of Advanced Research in Education*. – 2023. – Vol. 2, no. 4. – P. 1–5. – DOI: 10.56397/JARE.2023.07.02.
22. Chan C.K.Y., Colloton T. *Generative AI in higher education: The ChatGPT effect*. Routledge, 2024. – DOI: 10.4324/9781003459026.
23. Pireci Sejdiu N., Sejdiu S. The quiet transformation of higher education in the AI era // *Open Research Europe*. – 2025. – Vol. 5. – DOI: 10.12688/openreseurope.20715.1.
24. McDonald N., Johri A., Ali A., Hingle A. Generative artificial intelligence in higher education: Evidence from an analysis of institutional policies and guidelines // *arXiv*. – 2024. – DOI: 10.48550/arXiv.2402.01659.
25. Ayanwale M.A., Omeh C.B., Oyeniran F.M., Kanu C.C. Ethical compliance and institutional policy support for artificial intelligence integration in African TVET Education: A structural equation modeling approach // *PLoS One*. – 2025. – Vol. 20, no. 12. – Article no. e0335747. – DOI: 10.1371/journal.pone.0335747.
26. Jayasinghe S., Gamage K.A.A., Yang D., Cheng C., Disanayake C., Apeji U.D. Six institutional intervention areas to support ethical and effective student use of generative AI in higher education: A narrative review // *Education Sciences*. – 2026. – Vol. 16, no. 1. – Article no. 137. – DOI: 10.3390/educsci16010137.
27. Yang J., Xie W., Ni J. A framework for AI ethics literacy: development, validation, and its role in fostering students' self-rated learning competence // *Scientific Reports*. – 2025. – Vol. 15, no. 38030. – DOI: 10.1038/s41598-025-21977-5.
28. Askari M. Reliable but supervised: evaluating a generative AI-rubric model for consistent and fair assessment in postgraduate education // *Assessment & Evaluation in Higher Education*. – 2025. Published online 26 July. – DOI: 10.1080/02602938.2025.2537774.
29. Elzerman G. AI governance is a duty of care, not a branding exercise // *Times Higher Education Campus*. – 2026. 23 February. – URL: <https://www.timeshighereducation.com/campus/ai-governance-duty-care-not-branding-exercise> (дата обращения: 17.02.2026).
30. Абдусаламов Р.А. Основные направления правового регулирования использования искусственного интеллекта в сфере высшего образования // *Юридический вестник Дагестанского государственного университета*. – 2024. – Т. 52, № 4. – С. 135–143. – DOI: 10.21779/2224-0241-2024-52-4-135-143.
31. Ивахненко Е.Н., Никольский В.С. ChatGPT in higher education and science: a threat or a valuable resource? // *Высшее образование в России*. – 2023. – Т. 32, № 4. – С. 9–22. – DOI: 10.31992/0869-3617-2023-32-4-9-22.
32. Вerezubova N., Mindlin Yu., Sakovich N. Methodological aspects of application of generative artificial intelligence in pedagogical activities // *Современная наука: актуальные проблемы теории и практики. Серия: Гуманитарные науки*. – 2025. – № 8. – С. 74–78. – DOI: 10.37882/2223-2982.2025.08.09.
33. Ракитов А.И. Высшее образование и искусственный интеллект: эйфория и алармизм // *Высшее образование в России*. – 2018. – № 6. – С. 41–49. – URL: <https://vovr.elpub.ru/jour/article/view/1392> (дата обращения: 17.02.2026).
34. Войтов И.В. Перспективы правового регулирования искусственного интеллекта в высшем образовании в Российской Федерации // *Вестник Российской правовой академии*. – 2025. – № 5. – С. 161–172. – DOI: 10.33874/2072-9936-2025-0-5-161-172.
35. Лучшева Л. В. Социальные проблемы использования искусственного интеллекта в высшем образовании: задачи и перспективы // *Научный Татарстан*. – 2020. – № 4. – С. 84–89. – URL: <https://tatarica.org/ru/razdely/biblioteka-tatarika/nauchnyj-tatarstan/2020/socialnye-problemy-ispolzovaniya-iskusstvennogo-intellekta-v-vysshem-obrazovanii-zadachi-i-perspektivy> (дата обращения: 17.02.2026).
36. Вerezubova N.A., Yakovleva O.A., Kishkinova O.A. Этические и педагогические риски использования искусственного интеллекта в высшем образовании // *Международный журнал гуманитарных и естественных наук*. 2025. № 4-2 (103). С. 82–86. DOI: 10.24412/2500-1000-2025-4-2-82-86
37. Константинова Л.В., Ворожихин В.В., Петров А.М., Титова Е.С., Штыхно Д.А. Генеративный искусственный интеллект в образовании: дискуссии и прогнозы // *Открытое образование*. – 2023. – Т. 27, № 2. – С. 36–48. – DOI: 10.21686/1818-4243-2023-2-36-48.
38. Мухлаева Т.В. Генеративный искусственный интеллект: трансформации в образовании, перспективы и динамика // *Человек и обра-*

- зование. – 2025. – № 2. – С. 142–153. – DOI: 10.54884/1815-7041-2025-83-2-142-153.
39. Зазжигалкин А.В., Мансуров Т.Т., Мерецков О.В. Регулирование искусственного интеллекта в образовании // Компетентность. – 2024. – № 6. – С. 10–16. – URL: <https://www.asms.ru/news/vyshel-svezhiy-nomer-zhurnala-kompetentnost-6-2024.html> (дата обращения: 17.02.2026).
 40. Корнев А.А., Шинакова А.Д. Регулирование использования искусственного интеллекта в академической среде в России и за рубежом: вызовы и перспективы // Современная коммуникативистика. – 2026. – Т. 15, № 1. – С. 54–59. – DOI: 10.12737/2587-9103-2026-15-1-54-59.
 41. Ананин Д.П., Комаров Р.В., Реморенко И.М. «Когда честно – хорошо, для имитации – плохо»: стратегии использования генеративного искусственного интеллекта в российском вузе // Высшее образование в России. – 2025. – Т. 34, № 2. – С. 31–50. – DOI: 10.31992/0869-3617-2025-34-2-31-50.
 42. Stacey M., Mockler N. Analysing education policy: theory and method / ed. by M. Stacey, N. Mockler. London: Routledge, 2024. – 276 p. – DOI: 10.4324/9781003353379.
 43. Krippendorff K. Content analysis: an introduction to its methodology. 3rd ed. Los Angeles: SAGE, 2013. – 441 p. – ISBN: 978-1-4129-8319-0.
 44. Stracke C.M., Griffiths D., Pappa D., Bećirović S., Polz E. et al. Analysis of artificial intelligence policies for higher education in Europe // International Journal of Interactive Multimedia and Artificial Intelligence. – 2025. – Vol. 9. – P. 124–137. – DOI: 10.9781/ijimai.2025.02.011.
 45. Aristombayeva M., Satybaldiyeva R., Maung B.M., Lee D. Guiding the uncharted: the emerging (and missing) policies on Generative AI in higher education // Frontiers in Education. – 2025. – Vol. 10. – Article no. 1644081. – DOI: 10.3389/educ.2025.1644081.
 46. Xiao P., Chen Y., Bao W. Waiting, Banning, and Embracing: An Empirical Analysis of Adapting Policies for Generative AI in Higher Education // arXiv. – 2023. – DOI: 10.48550/arXiv.2305.18617.
 47. Rughiniş C., Vulpe S-N., Țurcanu D., Rughiniş R. AI at the knowledge gates: institutional policies and hybrid configurations in universities and publishers // Frontiers in Computer Science. – 2025. – Vol. 7. – Article no. 1608276. – DOI: 10.3389/fcomp.2025.1608276.
 48. Song Z., Kim C. The international status of generative artificial intelligence applications in higher education // Creative Computing. – 2025. – Vol. 6, no. 2. – P. 195–218. – DOI: 10.61131/cc.2025.6.2.195.
 49. Hofmann P., Späthe E., Brand A., Lins S., Sunyaev A. AI-based tools in higher education – A comparative analysis of university guidelines // Mensch und Computer 2024. ACM, 2024. – P. 1–6. – DOI: 10.1145/3670653.3677513.
 50. Perkins M., Roe J. Detection of GPT-4 generated text in higher education: Combining academic judgement and software to identify generative AI tool misuse // Journal of Academic Ethics. – 2023. – Vol. 22. – P. 89–113. – DOI: 10.1007/s10805-023-09492-6.
 51. Jin Y., Yan L., Echeverria V., Gašević D., Martinez-Maldonado R. Generative AI in higher education: A global perspective of institutional adoption policies and guidelines // Computers and Education: Artificial Intelligence. – 2025. – Vol. 8. – Article no. 100348. – DOI: 10.1016/j.caeai.2024.100348.

Статья поступила в редакцию 05.03.2026

Принята к публикации 10.04.2026

References

1. DiMaggio, P.J., Powell, W.W. (1983). The Iron Cage Revisited: Institutional Isomorphism and Collective Rationality in Organizational Fields. *American Sociological Review*. Vol. 48, no. 2, pp. 147–160. Available at: <http://www.jstor.org/stable/2095101> (accessed 17.02.2026).
2. Pratt, J. (2004). Institutional Isomorphism and Online learning in Australian Higher Education. In: *Inaugural Academy of World Business, Marketing and Management Development Conference. Academy of World Business, Marketing and Management Development*. Available at: <https://opus.lib.uts.edu.au/bitstream/10453/7529/1/2004000718.pdf> (accessed 17.02.2026).

3. Cardona Mejia, L.M., Pardo del Val, M., Dasi Coscollar, A. (2020). The Institutional Isomorphism in the Context of Organizational Changes in Higher Education Institutions. *International Journal of Research in Education and Science (IJRES)*. Vol. 6, no. 1, pp. 61-73, doi: 10.46328/ijres.v6i1.639.
4. Baharuddin, N.B., Abd Rahman, N.H., Kamil, N.L.M. (2025). Digital Education in Malaysian Public Universities: A Neo-Institutional Perspective on Policy Implementation. *International Journal of Modern Education*. Vol. 7, no. 27, pp. 571-588, doi: 10.35631/IJMOE.7270355.
5. Nemirovskiy, V.G. (2024). Artificial Intelligence in the Mirror of the Neo-Institutional Approach: A Sociological Analysis. *Sotsial'nye i gumanitarnye nauki. Otechestvennaya i zarubezhnaya literatura. Ser. 11, Sotsiologiya = Social Sciences and Humanities. Domestic and Foreign Literature. Series 11, Sociology*. No. 4, doi: 10.31249/rsoc/2024.04.03. (In Russ., abstract in Eng.)
6. Berger, P.L., Luckmann, T. (2020). *Sotsial'noe konstruirovaniye real'nosti: Dogovor o sotsiologii znaniya* [The Social Construction of Reality: A Treatise in the Sociology of Knowledge]. Transl. from English by A.A. Nazarenko. Moscow: Akademicheskii proekt Publ., 323 p. ISBN: 5-85691-036-2. (In Russ.)
7. Owidi, S.O., Lyanda, J.N., Okono, E.O., Micheni, E.M. (2025). The Dark Side of Generative AI in Higher Education: Challenges, Ethical Concerns, and Implications. *East African Journal of Arts and Social Sciences*. Vol. 8, no. 3, pp. 526-543, doi: 10.37284/eajass.8.3.3690.
8. Sharma, R.C., Panja, S.K. (2025). Addressing Academic Dishonesty in Higher Education: A Systematic Review of Generative AI's Impact. *Open Praxis*. Vol. 17, no. 2, pp. 251-269, doi: 10.55982/openpraxis.17.2.820.
9. AlBlooshi, S. (2026). Artificial Intelligence in Higher Education, Opportunities, and Challenges: A Review. *Frontiers in Education*. Vol. 10, article no. 1683968, doi: 10.3389/educ.2025.1683968.
10. Valentini, A., Blancas, A. (2025). The Challenges of AI in Higher Education and Institutional Responses: Is There Room for Competency Frameworks? *UNESCO IESALC Working Paper*. Available at: <https://www.iesalc.unesco.org/en/articles/challenges-ai-higher-education-and-imperative-competency-frameworks> (accessed 17.02.2026).
11. Jukiewicz, M. (2025). How Generative Artificial Intelligence Transforms Teaching and Influences Student Wellbeing in Future Education. *Frontiers in Education*. Vol. 10, article no. 1594572, doi: 10.3389/educ.2025.1594572.
12. Luo, J., Zheng, C., Yin, J., Teo, H.H. (2025). Design and Assessment of AI-based Learning Tools in Higher Education: A Systematic Review. *International Journal of Educational Technology in Higher Education*. Vol. 22, article no. 42, doi: 10.1186/s41239-025-00520-8.
13. Zawacki-Richter, O., Bai, J.Y.H., Lee, K., Slagter van Tryon, P.J., Prinsloo, P. (2024). New Advances in Artificial Intelligence Applications in Higher Education? *International Journal of Educational Technology in Higher Education*. Vol. 21, article no. 32, doi: 10.1186/s41239-024-00471-6.
14. Bond, M., Khosravi, H., De Laat, M., Bergdahl, N., Wegeia, V. et al. (2024). A Meta Systematic Review of Artificial Intelligence in Higher Education: A Call for Increased Ethics, Collaboration, and Rigour. *International Journal of Educational Technology in Higher Education*. Vol. 21, pp. 1-24, doi: 10.1186/s41239-024-00467-2.
15. Kalnina, D., Nimante, D., Baranova, S. (2024). Artificial Intelligence for Higher Education: Benefits and Challenges for Pre-Service Teachers. *Frontiers in Education*. Vol. 9, article no. 1501819, doi: 10.3389/educ.2024.1501819.
16. Baidoo-Anu, D., Owusu Ansah, L. (2023). Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learn-

- ing. *Journal of AI*. Vol. 7, no. 1, pp. 52-62, doi: 10.61969/jai.1337500.
17. Li, M., Xie, Q., Enkhtur, A., Meng, S., Chen, L. et al. (2025). A Framework for Developing University Policies on Generative AI Governance: A Cross-National Comparative Study. arXiv, doi: 10.48550/arXiv.2504.02636.
 18. Isaifan, R.J. (2026). Artificial Intelligence in Higher Education: A Global Statistical Synthesis for Policy and Quality Assurance Reform. *Education Sciences*. Vol. 16, no. 3, article no. 483, doi: 10.3390/educsci16030483.
 19. Karkoulilian, S., Sayegh, N., Sayegh, N. (2025). ChatGPT Unveiled: Understanding Perceptions of Academic Integrity in Higher Education – A Qualitative Approach. *Journal of Academic Ethics*. Vol. 23, pp. 1171-1188, doi: 10.1007/s10805-024-09543-6.
 20. Bernstein, S., Rahman, A., Sharifi, N., Terbish, A., MacNeil, S. (2025). Beyond the Benefits: A Systematic Review of the Harms and Consequences of Generative AI in Computing Education. arXiv. Available at: <https://arxiv.org/abs/2510.04443> (accessed 17.02.2026).
 21. Wu, Y. (2023). Integrating Generative AI in Education: How ChatGPT Brings Challenges for Future Learning and Teaching. *Journal of Advanced Research in Education*. Vol. 2, no. 4, pp. 1-5, doi: 10.56397/JARE.2023.07.02.
 22. Chan, C.K.Y., Colloton, T. (2024). *Generative AI in Higher Education: The ChatGPT Effect*. Routledge. DOI: 10.4324/9781003459026.
 23. Pireci Sejdiu, N., Sejdiu, S. (2025). The Quiet Transformation of Higher Education in the AI Era. *Open Research Europe*. Vol. 5, doi: 10.12688/openreseurope.20715.1.
 24. McDonald, N., Johri, A., Ali, A., Hingle, A. (2024). Generative Artificial Intelligence in Higher Education: Evidence from an Analysis of Institutional Policies and Guidelines. arXiv, doi: 10.48550/arXiv.2402.01659.
 25. Ayanwale, M.A., Omeh, C.B., Oyeniran, F.M., Kanu, C.C. (2025). Ethical Compliance and Institutional Policy Support for Artificial Intelligence Integration in African TVET Education: A Structural Equation Modeling Approach. *PLoS One*. Vol. 20, no. 12, article no. e0335747, doi: 10.1371/journal.pone.0335747.
 26. Jayasinghe, S., Gamage, K.A.A., Yang, D., Cheng, C., Disanayake, C., Apeji, U.D. (2026). Six Institutional Intervention Areas to Support Ethical and Effective Student Use of Generative AI in Higher Education: A Narrative Review. *Education Sciences*. Vol. 16, no. 1, article no. 137, doi: 10.3390/educsci16010137.
 27. Yang, J., Xie, W., Ni, J. (2025). A Framework for AI Ethics Literacy: Development, Validation, and Its Role In Fostering Students' Self-Rated Learning Competence. *Scientific Reports*. Vol. 15, article no. 38030, doi: 10.1038/s41598-025-21977-5.
 28. Askari, M. (2025). Reliable But Supervised: Evaluating a Generative AI-rubric Model for Consistent and Fair Assessment in Postgraduate Education. *Assessment & Evaluation in Higher Education*. Published online 26 July, doi: 10.1080/02602938.2025.2537774.
 29. Elzerman, G. (2026). AI Governance is a Duty of Care, Not a Branding Exercise. *Times Higher Education Campus*. 23 February. Available at: <https://www.timeshighereducation.com/campus/ai-governance-duty-care-not-branding-exercise> (accessed 17.02.2026).
 30. Abdusalamov, R.A. (2024). Main Directions of Legal Regulation of the Use of Artificial Intelligence in Higher Education. *Yuridicheskij vestnik Dagestanskogo gosudarstvennogo universiteta* = *Law Bulletin of Dagestan State University*. Vol. 52, no. 4, pp. 135-143, doi: 10.21779/2224-0241-2024-52-4-135-143 (In Russ., abstract in Eng.).
 31. Ivakhnenko, E.N., Nikol'skiy, V.S. (2023). ChatGPT in Higher Education and Science: A Threat or a Valuable Resource? *Vysshee obrazovanie v Rossii* = *Higher Education in Russia*. Vol. 32, no. 4, pp. 9-22, doi: 10.31992/0869-3617-2023-32-4-9-22 (In Russ., abstract in Eng.).

32. Verezubova, N., Mindlin, Yu., Sakovich, N. (2025). Methodological Aspects of Application of Generative Artificial Intelligence in Pedagogical Activities. *Sovremennaya nauka: aktual'nye problemy teorii i praktiki. Seriya: Gumanitarnye nauki = Contemporary Science: Current Theory and Practice. Series: Humanities*. No. 8, pp. 74-78, doi: 10.37882/2223-2982.2025.08.09 (In Russ., abstract in Eng.).
33. Rakitov, A.I. (2018). Higher Education and Artificial Intelligence: Euphoria and Alarmism. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. No. 6, pp. 41-49. Available at: <https://vovr.elpub.ru/jour/article/view/1392> (accessed 17.02.2026). (In Russ., abstract in Eng.).
34. Voytov, I.V. (2025). Prospects for Legal Regulation of Artificial Intelligence in Higher Education in the Russian Federation. *Vestnik Rossiyskoy pravovoy akademii = Bulletin of the Russian Law Academy*. No. 5, pp. 161-172, doi: 10.33874/2072-9936-2025-0-5-161-172 (In Russ., abstract in Eng.).
35. Luchsheva, L.V. (2020). Social Problems of Using Artificial Intelligence in Higher Education: Tasks and Prospects. *Nauchnyy Tatarstan*. No. 4, pp. 84-89. Available at: <https://tatarica.org/ru/razdely/biblioteka-tatarika/nauchnyj-tatarstan/2020/socialnye-problemy-ispolzovaniya-iskusstvennogo-intellekta-v-vyshhem-obrazovanii-zadachi-i-perspektivy> (accessed 17.02.2026). (In Russ., abstract in Eng.).
36. Verezubova, N.A., Yakovleva, O.A., Kishkinova, O.A. (2025). Ethical and Pedagogical Risks of Using Artificial Intelligence in Higher Education. *Mezhdunarodnyy zhurnal humanitarnykh i estestvennykh nauk = International Journal of Humanities and Natural Sciences*. No. 4-2 (103), pp. 82-86, doi: 10.24412/2500-1000-2025-4-2-82-86 (In Russ., abstract in Eng.).
37. Konstantinova, L.V., Vorozhikhin, V.V., Petrov, A.M., Titova, E.S., Shtykhn, D.A. (2023). Generative Artificial Intelligence in Education: Discussions and Forecasts. *Otkrytoe obrazovanie = Open Education*. Vol. 27, no. 2, pp. 36-48, doi: 10.21686/1818-4243-2023-2-36-48 (In Russ., abstract in Eng.).
38. Mukhlaeva, T.V. (2025). Generative Artificial Intelligence: Transformations in Education, Prospects and Dynamics. *Chelovek i obrazovanie = Man and Education*. No. 2, pp. 142-153, doi: 10.54884/1815-7041-2025-83-2-142-153 (In Russ., abstract in Eng.).
39. Zazhigalkin, A.V., Mansurov, T.T., Meretskov, O.V. (2024). Regulation of Artificial Intelligence in Education. *Kompetentnost' = Competence*. No. 6, pp. 10-16. Available at: <https://www.asms.ru/news/vyshel-svezhiy-nomer-zhurnala-kompetentnost-6-2024.html> (accessed 17.02.2026). (In Russ., abstract in Eng.).
40. Korenev, A.A., Shinakova, A.D. (2026). Regulation of the Use of Artificial Intelligence in the Academic Environment in Russia and Abroad: Challenges and Prospects. *Sovremennaya kommunikativistika = Modern Communication Studies*. Vol. 15, no. 1, pp. 54-59, doi: 10.12737/2587-9103-2026-15-1-54-59 (In Russ., abstract in Eng.).
41. Ananin, D.P., Komarov, R.V., Remorenko, I.M. (2025). "When Honesty is Good, for Imitation it is Bad": Strategies for Using Generative Artificial Intelligence in a Russian University. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 2, pp. 31-50, doi: 10.31992/0869-3617-2025-34-2-31-50 (In Russ., abstract in Eng.).
42. Stacey, M., Mockler, N. (2024). *Analysing Education Policy: Theory and Method*. Ed. by M. Stacey, N. Mockler. London: Routledge, 276 p., doi: 10.4324/9781003353379.
43. Krippendorff, K. (2013). *Content Analysis: An Introduction to Its Methodology*. 3rd ed. Los Angeles: SAGE, 441 p. ISBN: 978-1-4129-8319-0.
44. Stracke, C.M., Griffiths, D., Pappa, D., Bećirović, S., Polz, E. et al. (2025). Analysis of Artificial Intelligence Policies for Higher Education in Europe. *International Journal of Interactive Multimedia and Artificial Intelligence*. Vol. 9, pp. 124-137, doi: 10.9781/ijimai.2025.02.011.

45. Aristombayeva, M., Satybaldiyeva, R., Maung, B.M., Lee, D. (2025). Guiding the Uncharted: The Emerging (and Missing) Policies on Generative AI in Higher Education. *Frontiers in Education*. Vol. 10, article no. 1644081, doi: 10.3389/educ.2025.1644081.
46. Xiao, P., Chen, Y., Bao, W. (2023). Waiting, Banning, and Embracing: An Empirical Analysis of Adapting Policies for Generative AI in Higher Education. *arXiv*, doi: 10.48550/arXiv.2305.18617.
47. Rughiniş, C., Vulpe, S-N., Turcanu, D., Rughiniş, R. (2025). AI at the Knowledge Gates: Institutional Policies and Hybrid Configurations in Universities and Publishers. *Frontiers in Computer Science*. Vol. 7, article no. 1608276, doi: 10.3389/fcomp.2025.1608276.
48. Song, Z., Kim, C. (2025). The International Status of Generative Artificial Intelligence Applications in Higher Education. *Creative Computing*. Vol. 6, no. 2, pp. 195-218, doi: 10.61131/cc.2025.6.2.195.
49. Hofmann, P., Späthe, E., Brand, A., Lins, S., Sunyaev, A. (2024). AI-based Tools in Higher Education – A Comparative Analysis of University Guidelines. In: *Mensch und Computer 2024*. ACM, pp. 1-6, doi: 10.1145/3670653.3677513.
50. Perkins, M., Roe, J. (2023). Detection of GPT-4 Generated Text in Higher Education: Combining Academic Judgement and Software to Identify Generative AI Tool Misuse. *Journal of Academic Ethics*. Vol. 22, pp. 89-113, doi: 10.1007/s10805-023-09492-6.
51. Jin, Y., Yan, L., Echeverria, V., Gašević, D., Martinez-Maldonado, R. (2025). Generative AI in Higher Education: A Global Perspective of Institutional Adoption Policies and Guidelines. *Computers and Education: Artificial Intelligence*. Vol. 8, article no. 100348, doi: 10.1016/j.caeai.2024.100348.

*The paper was submitted 05.03.2026
Accepted for publication 10.04.2026*



Science Index РИНЦ-2024

Тематика «Народное образование»

ВОПРОСЫ ОБРАЗОВАНИЯ	9,277
ВЫСШЕЕ ОБРАЗОВАНИЕ В РОССИИ	8,803
ВЕСТНИК МЕЖДУНАРОДНЫХ ОРГАНИЗАЦИЙ: ОБРАЗОВАНИЕ, НАУКА, НОВАЯ ЭКОНОМИКА	8,332
ОБРАЗОВАНИЕ И НАУКА	7,955
ПСИХОЛОГИЧЕСКАЯ НАУКА И ОБРАЗОВАНИЕ	7,738
РУСИСТИКА	7,320
ЯЗЫК И КУЛЬТУРА	6,834
TRAINING, LANGUAGE AND CULTURE	6,773
ИНТЕГРАЦИЯ ОБРАЗОВАНИЯ	6,554
УНИВЕРСИТЕТСКОЕ УПРАВЛЕНИЕ: ПРАКТИКА И АНАЛИЗ	6,491

Вызовы и риски искусственного интеллекта для системы вузовского образования: понимание преподавателями и практическая готовность к изменениям

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-55-82

Сысоев Павел Викторович – д-р пед. наук, профессор, руководитель Научного центра Российской академии образования, SPIN-код: 2943-7230, ORCID: 0000-0001-7478-7828, Scopus Author ID: 8419258800, psysoyev@yandex.ru

Тамбовский государственный университет им. Г.Р. Державина, Тамбов, Россия

Адрес: 392000, г. Тамбов, Интернациональная ул., д. 33

***Аннотация.** Одним из основных противоречий интеграции технологий искусственного интеллекта (ИИ) в высшее образование является то, что, предоставляя беспрецедентные дидактические возможности для адаптации и персонализации обучения, ИИ одновременно порождает комплекс системных вызовов и рисков, требующих переосмысления содержания обучения, традиционных методов обучения и контроля, а также роли преподавателя в образовательном процессе. При этом эффективность противостояния возникающим угрозам во многом зависит от того, насколько преподаватели осознают данные вызовы и готовы практически на них ответить. Цель исследования – определить перечень ключевых вызовов и рисков, стоящих перед системой вузовского образования в условиях распространения ИИ, а также выявить степень понимания преподавателями высшей школы этих вызовов и рисков и их готовность к практическим действиям по их преодолению. На основе анализа научной литературы были выделены три основных вызова (изменение содержания обучения, изменение традиционных подходов и методов обучения и контроля, изменение роли преподавателя в триаде «преподаватель – ИИ – обучающийся») и четыре группы рисков (развитие клипового мышления и снижение когнитивной активности студентов, распространение ИИ-плагиата, «галлюцинации» и предвзятость ИИ, усиление цифрового и социального неравенства). Для определения понимания преподавателями этих вызовов и рисков, а также их готовности к ответу на них было проведено онлайн-анкетирование. В качестве респондентов выступили 1272 преподавателя из 43 вузов РФ. Результаты анкетного опроса свидетельствуют о наличии системного разрыва между декларируемым пониманием вызовов и рисков (около 70% педагогов осознают угрозы ИИ-плагиата и снижение когнитивных способностей студентов) и реальной практической готовностью к организационным и методическим изменениям (лишь 20–30% педагогов реально изменяют методы обучения и контроля, интегрируя ИИ в учебный процесс). Около 30–40% респондентов выразили нейтральное отношение по большинству вопросов, связанных с внедрением новых методов и форм контроля, а 15–25% опрошенных – отрицательное отношение. Получен-*

ные данные свидетельствуют о том, что на современном этапе системная интеграция ИИ в высшее образование не может осуществляться на инициативной основе отдельными преподавателями. Основная масса педагогов российских вузов не имеет системной поддержки (нормативной, методической, временной и финансовой) для перехода к новым моделям обучения в условиях распространения ИИ. Противостояние вызовам и рискам ИИ возможно исключительно на институциональном уровне при разработке и внедрении нормативных регламентов использования ИИ, введении универсальной компетенции в области ИИ во ФГОС ВО, пересмотре учебной нагрузки преподавателей и создании ресурсной базы для персонализированной работы студентов с инструментами ИИ.

Ключевые слова: искусственный интеллект, высшая школа, вызовы и риски интеграции искусственного интеллекта, ИИ-плагиат, цифровое и социальное неравенство

Для цитирования: Сысоев П.В. Вызовы и риски искусственного интеллекта для системы вузовского образования: понимание преподавателями и практическая готовность к изменениям // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 55–82. – DOI: 10.31992/0869-3617-2026-35-6-55-82.

Challenges and Risks of Artificial Intelligence for the Higher Education System: Faculty Understanding and Practical Readiness for Change

Original article

DOI: 10.31992/0869-3617-2026-35-6-55-82

Pavel V. Sysoyev – Dr.Sci. (Education), Professor, Director of the Research Center of the Russian Academy of Education, SPIN-code: 2943-7230, ORCID: 0000-0001-7478-7828, Scopus Author ID: 8419258800, psysoyev@yandex.ru

Derzhavin Tambov State University, Tambov, Russian Federation

Address: 33, Internatsyonalnaya str., Tambov, 392 000, Russian Federation

Abstract. One of the main contradictions in the integration of artificial intelligence (AI) technologies into higher education is that, while providing unprecedented didactic opportunities for adapting and personalizing learning, AI simultaneously creates a complex set of systemic challenges and risks that require a rethinking of educational content, traditional teaching and assessment methods, and the role of the teacher in the educational process. Moreover, the effectiveness of countering emerging threats largely depends on the extent to which faculty understand these challenges and are prepared to respond to them. The purpose of this study is to identify key challenges and risks facing the higher education system in the context of AI, as well as to determine the extent to which faculty understand these challenges and risks and their readiness to take practical steps to overcome them. Based on an analysis of the scientific literature, three key challenges were identified (changes in educational content, changes in traditional approaches and methods of teaching and assessment, and the changing role of the teacher in the “teacher-AI-student” triad) and four risk groups (the development of clip-based thinking and a decrease in students’ cognitive activity, the spread of AI plagiarism, hallucinations and bias in AI, and the strengthening of digital and social inequality). An online survey was conducted to determine lecturers’ understanding of these challenges and risks, as well as their preparedness to respond to them. Respondents included 1,272 lecturers from 43 Russian

universities. The survey results reveal a systemic gap between the stated understanding of challenges and risks (about 70% of teachers recognize the threat of AI plagiarism and the decline in students' cognitive abilities) and the actual practical readiness for organizational and methodological changes (only 20-30% of teachers are actually changing their teaching and monitoring methods, integrating AI into the educational process). Approximately 30-40% of respondents expressed a neutral attitude toward most issues related to the implementation of new methods and forms of assessment, while 15-25% expressed a negative attitude. The data obtained indicate that, at the current stage, the systemic integration of AI into higher education cannot be carried out on a proactive basis by individual faculty. The majority of Russian university faculty lack systemic support (regulatory, methodological, time-based, and financial) for the transition to new educational models in the context of the spread of AI. Addressing the challenges and risks of AI is possible only at the institutional level through the development and implementation of regulatory guidelines for the use of AI, the introduction of universal AI competencies in the Federal State Educational Standards of Higher Education, the revision of faculty workloads, and the creation of a resource base for personalized student engagement with AI tools.

Keywords: artificial intelligence, higher education, challenges and risks of artificial intelligence integration, AI plagiarism, digital and social inequality

Cite as: Sysoyev, P.V. (2026). Challenges and Risks of Artificial Intelligence for the Higher Education System: Faculty Understanding and Practical Readiness for Change. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 55-82, doi: 10.31992/0869-3617-2026-35-6-55-82 (In Russ., abstract in Eng.).

Введение

Появление и широкое распространение в системе образования искусственного интеллекта (ИИ) за последние несколько лет существенно видоизменило и разнообразило ландшафт научных исследований в области философии образования, педагогики, методики обучения дисциплинам, психологии и социологии. Учёные с позиции разных областей знания оценивают потенциал ИИ, его влияние на личность обучающегося и те изменения, которые он уже начинает постепенно привносить в обучение, развитие и воспитание учащихся и студентов. При этом следует специально отметить, что по некоторым позициям *система образования в целом оказалась не готова к такому стремительному вторжению ИИ*, который, предлагая решения для одних вопросов и проблем, *поднимает и обостряет другие, глубинные и более глобальные по своему масштабу и последствиям*.

На начальном этапе появления новых инновационных средств в центре внимания

учёных оказалось изучение дидактического и методического потенциала конкретных технологических решений на базе ИИ. Способность ИИ вступать в диалоговое взаимодействие с пользователем, по ряду факторов характеризующее общение между людьми, и предоставлять различные виды обратной связи [1] в адаптивном режиме существенно отличает его от других уже традиционных цифровых средств обучения. В результате в научной литературе наблюдается эффект «педагогического восторга», когда в центре внимания исследователей находятся преимущественно дидактические и методические возможности ИИ решать некоторые учебные, методические и исследовательские задачи. Безусловно, это объективно оправдано, когда учебное взаимодействие обучающихся с инструментами ИИ направлено на создание *дополнительных* возможностей глубже изучать профильные учебные дисциплины и формировать необходимые профессиональные компетенции. Примерами подобных исследований могут служить работы К. Чана

и Н. Зари [2], К.С. Итинсона [3], В. Жанга и соавторов [4], посвящённые использованию инструментов ИИ при подготовке студентов медицинских специальностей; А.Р. Ягудина, В.С. Цилицкого и соавторов [5], С. Феуерриегела и соавторов [6] – при подготовке будущих экономистов; А.П. Ивановой [7], П.В. Сысоева, М.В. Гаврилова, С.Ю. Булочникова [8] – юристов; А.Ф. Смыка и Е.А. Гусевой [9] – инженеров и т. п. Возможность генеративного ИИ к адаптации сложности текстов и разработке дидактических, учебных и контрольных материалов послужила стимулом создания содержательных моделей персонализированного обучения [10–12], реализуемых в том числе на базе интеллектуальных систем обучения (ИСО) [13].

Распространение ИИ в образовании, а также его значимые преимущества, заключающиеся в способности взять на себя решение многих задач, которые традиционно выступали функцией педагога, неизбежно приводят к *переосмыслению методологии образования* и началу дискуссии о *пересмотре некоторых положений традиционных подходов и методов обучения* [14; 15]. При этом авторы справедливо, на наш взгляд, аргументируют потребность в подобном переосмыслении *не только дидактическим и методическим потенциалом ИИ*, который уже получил освещение в научной литературе, но и, главное, *проблемными зонами*, которые по степени охвата и глубине распространения представляют реальные вызовы и угрозы для системы образования в целом и развития когнитивных способностей обучающихся в частности.

В своих работах учёные с разных методологических позиций раскрывают конкретные проблемы, стоящие перед системой образования в эпоху распространения ИИ. В частности, в центре внимания исследований Д. Коттона, П. Коттона, Дж. Шипвея [16] и П.В. Сысоева [17] находится проблема распространения в академической среде ИИ-плагиата – несанкционированного заимствования учащимися и студентами ма-

териалов обратной связи от генеративного ИИ; С.В. Титовой и И.В. Харламенко [18], П.В. Сысоева [19] – низкий по сравнению с обучающимися уровень сформированности у учителей и преподавателей компетенции в области ИИ; Я.И. Кузьмина, Е.В. Кручинской, И.А. Груздева, А.А. Наумова [20] – социальное разделение обучающихся по степени когнитивного развития ввиду широкой интеграции ИИ в образование и распространение ИИ-плагиата и др. Многие из приведённых и рассматриваемых отдельно вопросов имеют общие корни и представляют собой следствие более глубоких и масштабных проблем. Это обуславливает *необходимость комплексного изучения острых вызовов и рисков*, которые встают перед системой образования в связи с интеграцией ИИ.

Цель исследования – определить перечень ключевых вызовов и рисков, стоящих перед современной системой образования в условиях динамичной интеграции ИИ, выявить степень понимания преподавателями высшей школы данных вызовов и рисков и их готовности ответить на них.

Достижение данной цели конкретизируется четырьмя взаимосвязанными исследовательскими вопросами:

ИВ 1. Какие вызовы и риски для системы высшего образования порождает широкое распространение технологий ИИ на современном этапе?

ИВ 2. В какой мере преподаватели российских вузов осознают вызовы и риски, связанные с интеграцией ИИ в образовательных процесс?

ИВ 3. Насколько преподаватели российских вузов практически готовы к организационным и методическим изменениям в ответ на выявленные вызовы и риски ИИ?

ИВ 4. Существует ли системный разрыв между декларируемым пониманием преподавателями вызовов и рисков ИИ и их реальной готовностью к практическим ответным действиям, и в чём конкретно он проявляется?

Научная новизна исследования определяется совокупностью следующих позиций.

Во-первых, разработана структурированная классификация педагогических вызовов и рисков, порождаемых широким распространением ИИ в системе высшего образования. В отличие от существующих типологий, ограниченных, как правило, одним-двумя аспектами (академическая честность или трансформация профессий), предложенная классификация охватывает шесть взаимосвязанных направлений: изменение мира профессий и требований к специалистам, отражающееся в содержании обучения, трансформацию форм и методов обучения и контроля, деградацию когнитивных навыков и критического мышления обучающихся, распространение ИИ-плагиата, «галлюцинации» и идеологическую ангажированность больших языковых моделей, а также усиление цифрового и социального неравенства. Предложенная классификация нашла отражение в разделах анкеты для решения последующих ИВ.

Во-вторых, впервые на репрезентативной выборке преподавателей российских вузов ($N = 1272$, 43 университета различных типов и регионов) получены эмпирические данные, позволяющие одновременно и в сопоставимых показателях измерить как декларируемое понимание, так и практическую готовность к ответным действиям по каждому из выделенных направлений. Существующие отечественные исследования в данной области либо ограничены отдельными аспектами интеграции ИИ, либо проводились на локальных выборках одного-двух университетов, что не позволяло делать обобщения на уровне системы высшего образования в целом.

В-третьих, на основе полученных эмпирических данных введён и содержательно наполнен конструкт «адаптивного консерватизма» применительно к профессиональному поведению преподавателей российских вузов в условиях интеграции ИИ в образование. Конструкт фиксирует устойчивую диспозицию, при которой высокий уровень декларируемой осведомлённости о вызовах и

рисках ИИ сочетается с низкой готовностью к реальным изменениям в педагогической практике. В отличие от ранее описанных в международной литературе феноменов «технологического сопротивления» и «пассивной адаптации», «адаптивный консерватизм» характеризует не отрицание необходимости изменений, а систематическое откладывание их практической реализации при внешнем принятии самой идеи трансформации.

В-четвёртых, предложена авторская шкала допустимого использования ИИ в учебной и исследовательской работе, обеспечивающая разграничение между санкционированным применением инструментов ИИ и ИИ-плагиатом и применимая для использования на уровне вуза или кафедры.

Обзор литературы

В настоящее время в научной литературе только начинается дискуссия о вызовах и рисках, которые угрожают человечеству в эпоху динамичного развития и распространения ИИ. Под *вызовом* понимается внешнее провокационное воздействие, требующее реакции, которое при игнорировании может перерасти в угрозу. *Риск* – это *сочетание вероятности возникновения опасных и нежелательных последствий в результате вызова*. В зависимости от направленности угрозы риски широкого распространения ИИ могут быть *философскими*, когда влияние ИИ на человека может способствовать изменениям в его мировоззрении; *психологическими*, когда постоянное общение с голосовыми помощниками и чат-ботами на базе ИИ может привести к психологическим расстройствам у человека (например, асоциализации и искажённому восприятию реальности); *относящимися к сфере информационной безопасности*, когда при развёртывании систем ИИ может произойти утечка персональных данных пользователей; *педагогическими*, когда интеграция ИИ в образование может привести к проблемам, связанным с процессом и результатом приобретения новых зна-

ний и развития умений и т. п. Отметим, что в реальной жизни многие вызовы и риски, а также их последствия тесно переплетаются друг с другом, усиливаясь или ослабевая под воздействием разных факторов. В центре внимания данного исследования находятся педагогические вызовы и риски, стоящие перед системой образования.

Педагогические вызовы к системе образования

Интеграция ИИ послужила причиной появления новых вызовов, с которыми столкнулась современная система высшего образования. К наиболее масштабным из них можно отнести: а) изменение содержания обучения; б) изменение традиционных подходов и методов обучения и контроля; в) изменение роли преподавателя в триаде «преподаватель – искусственный интеллект – обучающийся». Рассмотрим подробнее каждый из них.

Изменение содержания обучения

Интеграция ИИ в образование будет способствовать ряду изменений в содержании обучения. *Во-первых*, в условиях распространения в молодёжной среде клипового мышления особую актуальность приобретает *переход от «знаниевой» к «деятельностной» модели обучения*, что было отражено А.Н. Леонтьевым и С.А. Рубинштейном в положениях деятельностного подхода и что, к сожалению, не всегда находит системное воплощение на практике. В ситуациях, когда ИИ может мгновенно предоставить необходимую фактологическую информацию, ценность её простого получения и запоминания падает. Обучение должно перейти от передачи знаний к обучению работе с постоянно изменяющимися массивами информации. На передний план выходят практические умения использовать знания в реальной деятельности, подвергать анализу и критическому осмыслению материалы обратной связи от ИИ, работать в командах, находить нестандартные решения, встраивать новые приобретённые знания и опыт в имеющуюся у обучающихся систему знаний.

Одним из ярких примеров доминирования деятельностной модели обучения выступает методика обучения иностранным языкам. Для овладения языком *как* средством общения его необходимо изучать *через* непосредственное общение. В результате такого обучения учащиеся и студенты не просто заучивают новую лексику и знакомятся в теории с грамматическими конструкциями, а развивают необходимые умения иноязычной устной и письменной речевой деятельности.

Во-вторых, учитывая дидактический потенциал ИИ, *компетенция в области ИИ* должна войти в перечень универсальных компетенций и стать частью содержания обучения студентов всех направлений подготовки и специальностей. Данный вид компетенции у студентов вузов может включать в себя два модуля: инвариантный и вариативный. *Инвариантный модуль* направлен на формирование способностей обучающихся корректно взаимодействовать с инструментами генеративного ИИ для решения учебных и исследовательских задач, правильно формулировать запросы, анализировать и подвергать критическому осмыслению материалы обратной связи от ИИ. *Вариативный модуль* ориентирован на использование конкретных инструментов ИИ для решения профессиональных задач. В методической литературе имеются исследования, посвящённые использованию технологических решений на базе ИИ в профессиональной подготовке студентов разных направлений подготовки и специальностей: медиков [3–5], юристов [6; 7], экономистов [2; 8], лингвистов [21], инженеров [9] и т. д. Овладение студентами профессиональными инструментами ИИ для решения профессиональных задач будет направлено на развитие у них учебной автономии, а также способствовать их востребованности на современном рынке труда.

В-третьих, возможность технологических решений на базе ИИ брать на себя решение многих второстепенных задач в различных профессиональных сферах, а также их постоянное совершенствование неизбеж-

но приведут к *значительным изменениям в мире профессий*. По причине объективного сокращения финансовых затрат ряд профессий полностью исчезнет. Требования к специалистам многих других сфер видоизменяются. Появятся новые профессии для решения новых задач в новом мире с ИИ. Университеты должны гибко реагировать на потребности секторов экономики, по сути, готовя специалистов «завтрашнего дня». Для этого потребуются с учётом способности предвидения разрабатывать основные профессиональные образовательные программы (ОПОП) нового содержания, изменять содержание учебных программ у существующих ОПОП, объединять или закрывать некоторые специальности.

Изменение традиционных подходов и методов обучения и контроля

Широкое применение технологий ИИ на практике привело к актуализации двух глубоких проблем. Первая связана с *девальвацией ценности фундаментального образования и научного мышления*. На протяжении столетий образование характеризовалось фундаментальностью и системностью, ориентированностью на глубокое изучение теоретических дисциплин, принципов науки и закономерностей, связью гуманитарного и естественнонаучного знаний, долгосрочностью и направленностью на развитие аналитических способностей, понимание причинно-следственной связи между фактами, событиями и явлениями, критического и научного мышления. В процессе обучения учащиеся и студенты развивали умения учебной автономии, позволяющие в дальнейшем саморазвиваться в парадигме «обучение на протяжении всей жизни» с целью удовлетворения личностных или профессиональных интересов и потребностей. Новые знания надстраивались на прочном каркасе имеющихся знаний и опыта практической деятельности. В идеале в этом контексте ценности фундаментального образования появление ИИ позволило бы обучающимся делегировать инструментам ИИ решение

ряда второстепенных задач или операций и в большей мере сфокусироваться на решении других, характеризующихся высоким когнитивным уровнем сложности задач. В реальности стремительное распространение ИИ с возможностью быстрой генерации готовых ответов способствует девальвации ценности фундаментального образования и научного мышления, *стимулирует формирование клипового мышления и распространение в академической среде ИИ-плагиата*. Не видя ценности в формировании системных знаний, многие обучающиеся предпочитают получать быстрые и не всегда корректные ответы от генеративного ИИ без глубокого понимания природы этой информации.

Вторая причина связана с *использованием дидактического потенциала ИИ в учебном процессе*. Возможность инструментов ИИ взять на себя решение ряда второстепенных задач требует от педагогов изменений в традиционных формах и методах обучения с тем, чтобы практика студентов с ИИ расширяла их возможности получения дополнительных знаний или развития необходимых умений.

На современном этапе распространения ИИ изменения в подходах и методах обучения должны коснуться интеграции внеаудиторной практики взаимодействия учащихся и студентов с конкретными технологическими решениями на базе ИИ в традиционные методики обучения дисциплинам. По предметно-тематическому содержанию и уровню сложности данная практика должна соотноситься с программными требованиями. Материалы внеаудиторного взаимодействия студентов с ИИ должны выступать предметом обсуждения на последующих аудиторных занятиях. Обучающиеся должны предоставить «цифровой след» работы с ИИ (скриншоты диалога с ИИ, историю создания версий письменной работы и т. п.). Такой подход к обучению будет способствовать формированию у обучающихся критического мышления, аналитических навыков и ответственности за процесс и результат

обучения. Похожие изменения затрагивают использование методов контроля. Специализированные инструменты ИИ способны разрабатывать задания и осуществлять автоматизированный контроль знаний обучающихся. Возможность средств генеративного ИИ предоставлять оценочную и корректирующую обратную связь создаёт необходимые условия для автоматизированной проверки творческих письменных работ (например, эссе, рефератов, курсовых и квалификационных работ, научных статей). В научной литературе описывается эффективность включения этапов взаимодействия обучающихся с ИИ в традиционные методики преподавания профильных дисциплин [21–24].

Безусловно, органичная интеграция ИИ в учебный процесс и изменения в традиционных подходах и методах обучения потребуют со стороны преподавателя пересмотра отношения к ИИ как к «участнику» учебного процесса в триаде «преподаватель – искусственный интеллект – обучающийся». Способность ИИ к диалоговому взаимодействию с обучающимися и преподавателями по адаптивности и персонализации значительно отличает его от других ранее известных средств обучения, предоставляющих в большей мере шаблонные ответы. В этой связи представляется своевременным и важным начало дискуссии относительно статуса ИИ в учебном процессе. И если пока ещё не пришло время считать ИИ полноценным субъектом образовательного процесса, то статус социального актора – «коммуникативного искусственного интеллекта» в рамках коммуникативного конструктивизма, по мнению В.С. Никольского [25], представляется весьма оправданным и определяет необходимость выделения специфической предметной области исследований в философии образования для его изучения.

Вместе с тем, несмотря на возможные изменения в традиционных подходах и методах обучения и контроля, позволяющих

интегрировать практику обучающихся с ИИ в традиционные методики обучения, пока ещё нерешёнными остаются вопросы оценивания *продуктов сотворчества ИИ и человека*. Как оценить работу, 5, 10, 20 % и т. п. объёма которой обучающийся выполнил с опорой на ИИ? Как отделить, например, использование ИИ для языкового редактирования текста эссе от переписывания работы по рекомендациям ИИ? Возможным решением станет принятие регламента с обозначением сфер и степени использования ИИ при выполнении учебных или исследовательских задач по определённой шкале допустимого использования (см., например, *Табл. 1*).

Окончательное решение о критериях допустимого использования ИИ должны принимать специалисты в каждой конкретной сфере. С целью предотвращения субъективности педагогов при использовании данных критериев представляется необходимым разработка локальных (кафедральных или факультетских) документов, содержащих уточнённые требования и примеры легального и допустимого использования ИИ в учебной и исследовательской работе.

Широкое распространение ИИ-плагиата (см. ниже) в академической среде обуславливает необходимость пересмотра некоторых форматов заданий или форм итогового контроля. Учитывая масштабное распространение ИИ-плагиата, представляется целесообразным пересмотреть отдельные форматы проведения государственной итоговой аттестации (ГИА). В частности, требуется либо существенная трансформация содержания выпускной квалификационной работы, либо её замена иными формами проведения ГИА. К числу потенциально эффективных решений можно отнести: устный экзамен, включающий критический анализ решений, сгенерированных ИИ, творческие проекты, практико-ориентированные кейсы, эссе, отражающие личный опыт и ценностные рассуждения обучающихся.

Таблица 1

Шкала допустимого использования ИИ в учебной и исследовательской работе

Table 1

Scale for acceptable use of AI in educational and research work

Уровень	Допустимое использование	Примеры	Действия студентов	Подход к оцениванию работы
Зелёный (разрешён без ограничений)	ИИ как инструмент поддержки	Редактирование текста, поиск литературы, генерация идей и аргументов, статистическая обработка данных	Указать, какие технологические решения на базе ИИ использовались и для решения каких задач; приложить запросы к ИИ	Работа оценивается по обычным критериям
Жёлтый (разрешено с частичными ограничениями)	ИИ для выполнения части работы	Генерация черновой версии, перевод текста, создание примеров	Указать, какие технологические решения на базе ИИ использовались и для решения каких задач; приложить запросы к ИИ; обязательно указать, что сделано с помощью ИИ, а что лично автором	Оценка снижается, если доля работы, созданной с использованием ИИ, больше 50%
Красный (полностью запрещено)	Исключительно самостоятельное выполнение работы	Устный экзамен, защита проекта, написание рефлексивного эссе о личном опыте	Не использовать ИИ	При любом использовании ИИ работа не засчитывается

Изменение роли преподавателя в триаде «преподаватель – искусственный интеллект – обучающийся»

Социальный заказ на педагогов, способных органично использовать ИИ для поднятия учебного процесса на более высокий по степени решения когнитивных задач уровень, представляет определённый вызов для системы высшего образования. С одной стороны, преподаватель сохраняет все имеющиеся функции по организации учебного процесса и обучения студентов. Делегируя инструментам ИИ свои второстепенные функции (проверку некоторых видов домашнего задания, разработку планов занятий, тренировочных упражнений и творческих заданий), педагог продолжает осуществлять контроль и критически оценивать материалы обратной связи от ИИ. С другой – в условиях интеграции ИИ в образование преподаватель должен овладеть компетенцией в области ИИ, позволяющей ему успешно реализовать ряд новых функ-

ций, связанных с организацией учебного процесса с использованием ИИ и обучением студентов использовать средства ИИ для разумной реализации моделей адаптивного или персонализированного обучения [10–12]. В последние годы появилось несколько работ, в которых учёными предлагалось содержание компетенции педагога в области ИИ [18; 19]. Анализ исследований показывает, что ключевыми компонентами данного вида компетенции должны быть следующие: а) мотивационно-целевой; б) нормативный правовой; в) информационная безопасность; г) этический; д) промпт-инжиниринг; е) обучение и контроль; ж) управление учебным процессом; з) профессиональное развитие [19].

Предлагаемое содержание компонентов компетенции преподавателей в области ИИ может изменяться с разработкой новых технологических решений на базе ИИ и/или разработкой новых методик обучения на базе ИИ. Одним из требований к профессиональ-

ной деятельности преподавателя в условиях интеграции ИИ выступает готовность к организационно-методическим изменениям, направленным на создание благоприятных условий для обучения и развития обучающихся.

Следует специально подчеркнуть, что несмотря на весь имеющийся у инструментов ИИ потенциал, его интеграция в систему образования лишь усиливает нагрузку на *педагога*, у которого появляются *дополнительные функции, требования к квалификации и ответственность*. По объективным причинам эти дополнительные функции увеличивают трудозатраты при том же нормировании рабочего времени преподавателя, что *может способствовать профессиональному выгоранию педагогов*, формальному отношению к новым требованиям, отказу от внедрения ИИ и текучке кадров. Администраторам системы образования необходимо принимать во внимание данный риск и задуматься над пересмотром учебной нагрузки и/или введением дополнительных ставок по сопровождению ИИ-блока.

Представленные выше вызовы имеют отношение ко всей системе высшего образования. При этом представители педагогического сообщества должны понимать, что неспособность вовремя и адекватно отреагировать на данные вызовы приведёт к определённым *рискам* – нежелательным последствиям в результате интеграции ИИ в образование. На настоящий момент к наиболее актуальным педагогическим рискам можно отнести следующие: а) развитие клипового мышления и снижение когнитивной активности студентов; б) распространение ИИ-плагиата в академической среде; в) способность ИИ к «галлюцинациям», предвзятость и идеологическая ангажированность и г) усиление цифрового и социального неравенства. Рассмотрим подробнее содержание каждого из них.

Развитие клипового мышления и снижение когнитивной активности студентов

Постоянно обновляющийся поток информации, доступный современной молодё-

жи посредством социальных сетей, породил *клиповое мышление*, когда учащиеся и студенты привыкают и предпочитают воспринимать информацию короткими сегментами. В результате обучающиеся усваивают получаемую информацию дискретными фрагментами, которые очень часто не образуют семантические связи с уже имеющимися у человека знаниями. Отсутствие глубокого анализа не позволяет молодым людям создать целостную картину, включающую наличие причинно-следственной связи между фактами, событиями и явлениями. Это, в свою очередь, не позволяет обучающимся развивать критическое мышление. Распространение генеративного ИИ значительно усугубило развитие клипового мышления. Не осознавая ценности фундаментального образования и развития научного мышления, как было отмечено выше, многие учащиеся и студенты вместо глубокого изучения дисциплины и единоличного выполнения заданий нередко предпочитают получать готовые ответы от генеративного ИИ. В перспективе это может привести к атрофии навыков самостоятельного формулирования проблем, выдвижения гипотез, разделения цели на соответствующие более детальные задачи и их решения в процессе самостоятельной мыслительной деятельности. Регулярное обращение к инструментам ИИ, предоставляющим готовые решения на запросы пользователя, создаёт у студентов *иллюзию эффективности при минимальном использовании собственных когнитивных усилий*. Их роль в учебном процессе из активных производителей или добытчиков знаний переходит в пользователей и редакторов запросов к ИИ. Вместо того, чтобы использовать имеющийся у средств ИИ дидактического и методического потенциала для выполнения шаблонных и рутинных функций с целью освобождения дополнительного времени для решения более сложных по уровню выполнения когнитивных задач, обучающиеся нередко предпочитают нарушать правила академической этики и допускать

ИИ-плагиат (см. ниже). Неспособность системы образования своевременно и гибко реагировать на данный вызов постепенно приводит к *снижению когнитивной способности обучающихся и их когнитивной деградации*. Вероятным способом противодействия указанной угрозе может выступать изменение подходов к обучению и контролю, при которых: а) внеаудиторная практика взаимодействия студентов с инструментами ИИ интегрируется в традиционную методику обучения; б) обучающиеся осуществляют критический анализ материалов обратной связи от генеративного ИИ; в) на аудиторных занятиях студенты должны предоставить «цифровой след» взаимодействия с ИИ; г) использование ИИ в образовательном процессе базируется на принципах академической этики и регламентируется локальными нормативными актами.

Распространение ИИ-плагиата в академической среде

Девальвация ценностей фундаментального образования, клиповое мышление и генеративные возможности ИИ послужили катализатором распространения в академической среде ИИ-плагиата. Речь идёт о генерации фрагментов или полных текстов курсовых и квалификационных работ. Угроза данного явления заключается в том, что, во-первых, как отмечают Д. Коттон, П. Коттон и Дж. Шипвей [16], Т. Волтцер и А. Дахл [26], ИИ-плагиат носит масштабный характер. Во-вторых, как показывает в своём эмпирическом исследовании П.В. Сысоев [17], порядка половины опрошенных студентов российских вузов воспринимают заимствование текста, сгенерированного ИИ, в качестве нормы. Они искренне распространяют авторство на материалы обратной связи от ИИ, сгенерированные по их индивидуальным запросам (промптам). Решений данной проблемы может быть несколько: от изменения традиционных подходов и методов обучения и контроля, включающих задачи и задания, выполняемые исключительно человеком, до формирования компетенции

преподавателей в области ИИ (их способности объяснить учащимся и студентам правила академической этики) и наличия в образовательных учреждениях локальных актов, регламентирующих сферу и степень использования средств генеративного ИИ в обучении, а также меры ответственности за их нарушение.

Способность ИИ к «галлюцинациям», предвзятость и идеологическая ангажированность

Важным риском интеграции ИИ в образование выступает неудовлетворительное качество материалов обратной связи от генеративного ИИ. Во-первых, способность ИИ к «галлюцинациям» позволяет средствам генеративного ИИ в случаях ограниченного доступа к необходимым базам данных при составлении ответов на запросы предлагать пользователям ложные сведения или вымышленную информацию. В научной литературе имеются многочисленные свидетельства такого неудовлетворительного решения поставленных перед ИИ задач, что ставит под вопрос использование ИИ при подготовке исследовательских работ [27]. Во-вторых, каждый инструмент генеративного ИИ функционирует на определённой большой языковой модели (БЯМ), которая отражает идеологию и социально-культурные традиции конкретной страны-разработчика. В этой связи сгенерированные материалы могут отличаться определённой предвзятостью в трактовке конкретных событий или фактов и иметь идеологическую ангажированность [28–31]. Примерами различий в фундаментальных установках может служить полярно разная интерпретация причин, событий и финала русской революции и Гражданской войны 1917–1922 гг., Великой Отечественной войны, холодной войны, окончания Второй мировой войны на Тихом океане (бомбардировка Хиросимы и Нагасаки в августе 1945 г.), войны за независимость США (1775–1783 гг.). Подобных примеров можно привести множество. Всё это подрывает доверие к генеративным ин-

струментам ИИ и будет требовать усилий со стороны пользователя, чтобы постоянно подвергать получаемую информацию перепроверке, критическому осмыслению и переработке, что значительно ограничивает сферу использования ИИ.

Усиление цифрового и социального неравенства

Традиционно развитие и распространение информационных и коммуникационных технологий и технологий искусственного интеллекта принято связывать со снижением социального неравенства в сфере образования и возможностью предоставления равного доступа к образовательным ресурсам (открытым онлайн-курсам или информационно-справочным материалам) обучающимся независимо от географического места проживания, возраста, социально-экономического класса, физических способностей и т. п. Это положение позиционируется во многих рекомендациях ЮНЕСКО¹. На первый взгляд интеграция ИИ в образование, представляющая собой следующий эволюционный шаг после распространения информационных и коммуникационных технологий, должна продолжить эстафету обеспечения всех пользователей сети Интернет равным доступом к возможностям ИИ в образовательных целях. Вместе с тем ряд факторов объективно определяют усиление цифрового и социального неравенства среди обучающихся. Во-первых, многие разработчики программного обеспечения (ПО) на базе ИИ предоставляют пользователям бесплатный доступ исключительно к демоверсиям программ или ограниченным по функциональным возможностям версиям ПО. Полноценный доступ к продуктам потребует финансовых затрат. Учитывая узкую профессиональную направленность большинства технологических решений на базе ИИ, доступ к ПО на базе ИИ для ре-

шения широкого спектра учебных или профессиональных задач будет определяться экономическим классом пользователя.

Во-вторых, в более глобальном плане, как справедливо отмечают в своей работе Я.И. Кузьминов, Е.В. Кручинская, И.А. Груздев и А.А. Наумов [20], социальное неравенство может и будет происходить на уровне учебных заведений по различиям в институциональной культуре работы с ИИ. Учреждения, которые по объективным причинам ориентированы на обучение студентов с более высоким уровнем общекультурного и когнитивного развития, мотивацией в получении образования и амбициозными планами на будущее в выбранной профессии, станут более ответственно подходить к использованию ИИ в образовании. В таких вузах уже разрабатываются и внедряются нормативные и правовые документы (регламенты), определяющие сферу и степень использования конкретных технологических решений на базе генеративного ИИ в профессиональной подготовке студентов [21; 32]. Отметим, что это системная работа в вузах, которая требует решения целого ряда вопросов: от повышения компетенции преподавателей и студентов в области ИИ до пересмотра и корректировки некоторых методов и форм обучения. В российской научной литературе имеются работы, описывающие опыт участия отечественных вузов в подобных изменениях [18; 33–35]. Такой подход позволит поднять учебный процесс на более высокий по степени решения когнитивных задач уровень. В других же учебных заведениях, где руководство выбирает позицию «выжидания» и продолжает официально запрещать или игнорировать использование студентами ИИ, будет процветать ИИ-плагиат и продолжится постепенная когнитивная деградация обучающихся. Таким образом, в данном ракурсе интеграция

¹ UNESCO. Guidance for generative AI in education and research. Paris, UNESCO, 2023. // Сайт ЮНЕСКО. – URL: <https://unesdoc.unesco.org/ark:/48223/pf0000386693> (дата обращения: 20.04.2026); UNESCO. AI and education: guidance for policy-makers. Paris, UNESCO, 2022. // Сайт ЮНЕСКО. – URL: <https://unesdoc.unesco.org/ark:/48223/pf0000376709> (дата обращения: 20.04.2026).

ИИ может способствовать усилению цифрового и социального неравенства.

Между выявленными вызовами и рисками для системы образования существует тесная взаимосвязь, что обуславливает необходимость их системного, а не дискретного решения.

Материалы и методы

В исследовании приняли участие 1272 преподавателя из 43 российских вузов, включая Московский государственный университет имени М.В. Ломоносова (МГУ), Московский городской педагогический университет (МГПУ), Московский педагогический государственный университет (МПГУ), Национальный исследовательский университет «Высшая школа экономики» (НИУ ВШЭ), Санкт-Петербургский политехнический университет Петра Великого (СПбПУ), Санкт-Петербургский государственный университет (СПбГУ), Нижегородский государственный лингвистический университет имени Н.А. Добролюбова (НГЛУ), Казанский (Приволжский) федеральный университет (КФУ), Белгородский государственный национальный исследовательский университет (НИУ БелГУ), Ивановский государственный университет (ИвГУ), Воронежский государственный университет (ВГУ), Воронежский государственный педагогический университет (ВГПУ), Воронежский государственный аграрный университет имени императора Петра I (ВГАУ), Мичуринский государственный аграрный университет (МичГАУ), Российскую академию народного хозяйства и государственной службы при Президенте Российской Федерации (РАНХиГС), Национальный исследовательский Томский государственный университет (НИ ТГУ), Томский политехнический университет (ТПУ), Тамбовский государственный университет имени Г. Р. Державина (ТГУ им. Г.Р. Державина) и др. Выбор учебных заведений, среди которых национальные исследовательские и региональные вузы, определялся желанием преподавателей принять участие в он-

лайн-анкетировании на платформе Яндекс.Формы. Согласно материалам, возрастной диапазон респондентов составил от 24 до 63 лет (24–29 лет ($n = 381$), 30–39 лет ($n = 357$), 40–50 лет ($n = 285$), 50–63 лет ($n = 249$)). Преподаватели читали профильные дисциплины по следующим направлениям подготовки/специальностям: «Педагогическое образование» (10,6%), «Экономика» (9,9%), «Юриспруденция» (9,6%), «Лингвистика» (8,3%), «Журналистика» (7,4%), «Филология» (7,3%), «История» (7%), «Лечебное дело» (7%), «Социология» (6,4%), «Прикладная информатика» (5,8%), «Информационная безопасность» (5,1%), «Биология» (4,2%), «Инфокоммуникационные технологии и системы связи» (3,9%) и др.

В качестве инструмента для определения понимания преподавателями высшей школы вызовов и рисков развития системы вузовского образования в условиях распространения ИИ и их готовности ответить на данные вызовы и риски выступила анкета, состоящая из восьми блоков, соответствующих обозначенным в данной работе вызовам и рискам. Под «готовностью» ответить на риски мы понимаем способность (намерение и имеющийся практический опыт) преподавателей гибко изменять подходы и методы обучения, формы осуществления контроля и взаимодействовать со средствами генеративного ИИ в соответствии с обозначенными угрозами. Респондентам было предложено на анонимной основе выразить своё отношение по каждому утверждению по симметричной биполярной шкале Ликерта (-2 – «полностью не согласен», -1 – «не согласен», 3 – «нейтральное отношение», 4 – «согласен», 5 – «полностью согласен»).

Результаты исследования

Результаты анкетирования по определению понимания преподавателями вызовов и рисков развития системы вузовского образования в условиях распространения ИИ и их готовности ответить на данные вызовы и риски представлены в *таблице 2*.

Таблица 2

Результаты анкетирования преподавателей высшей школы на предмет понимания вызовов и рисков развития системы вузовского образования в условиях распространения ИИ и их готовности ответить на данные вызовы и риски

Table 2

Results of a survey of higher education teachers on their understanding of the challenges and risks of developing the higher education system in the context of the spread of AI and their readiness to respond to these challenges and risks

Утверждение	Варианты ответа (%)					Статистические характеристики	
	-2	-1	0	1	2	Среднее $\bar{x}$	Мода (M_0)
1. Вызов: Изменение содержания обучения							
1.1. Я осознаю, что с интеграцией ИИ многие профессии видоизменяются. Я готов(-а) адаптировать содержание читаемых мною дисциплин под новые профессиональные задачи	1,9	5,0	8,2	30,8	54,1	1,27	2
1.2. Я готов(-а) пересмотреть содержание учебных дисциплин, изменив акцент с передачи знаний на развитие умений использовать знания на практике	11,9	14,4	20,1	32,7	20,9	0,34	1
1.3. Я считаю необходимым включить компетенцию в области ИИ в число УК студентов всех направлений подготовки или специальностей	3,1	7,5	14,5	29,0	45,9	1,05	2
1.4. Я разрабатываю задания, в которых студенты используют ИИ для анализа данных, а не для простого воспроизведения информации	10,9	25,7	12,9	34,0	16,5	0,20	1
1.5. Я вижу необходимость в регулярном обновлении ОПОП из-за появления новых профессиональных инструментов на базе ИИ	3,2	17,5	35,7	30,8	12,8	0,32	0
1.6. Я готов(-а) оценивать не столько знания студентов, сколько их способность применять знания на практике с использованием ИИ или без него	2,6	16,9	33,2	31,4	15,9	0,41	1
2. Вызов: Изменение традиционных подходов и методов обучения и контроля							
2.1. Я считаю, что учебное взаимодействие студентов с инструментами ИИ должно осуществляться во внеаудиторное время	5,6	11,9	20,8	29,6	32,1	0,71	2
2.2. Я организую внеаудиторную практику студентов с инструментами ИИ и интегрирую её в традиционную методику преподавания дисциплины	3,3	27,5	18,2	33,3	17,7	0,35	1
2.3. Студенты обязательно должны предоставлять «цифровой след» взаимодействия с ИИ (скриншоты диалога, историю версий работы и т. п.)	8,1	13,8	20,1	27,0	30,8	0,59	2
2.4. Я пересматриваю традиционные формы контроля в связи с использованием студентами инструментов ИИ при выполнении домашних учебных заданий и исследовательских работ	13,2	28,2	23,8	20,8	14,0	-0,07	1
2.5. Я объясняю студентам разницу между допустимым и недопустимым использованием ИИ в учебном процессе	1,9	15,0	26,9	37,1	19,1	0,62	1
2.6. Я разрабатываю задания, которые невозможно выполнить с помощью ИИ без глубокого понимания материала	3,8	18,8	27,6	33,3	16,5	0,39	1

Продолжение таблицы 2

Утверждение	Варианты ответа (%)					Статистические характеристики	
	-2	-1	0	1	2	Среднее $\bar{x}$	Мода (M_o)
2.7. Я вижу необходимость в разработке локальных (кафедральных, факультетских) документов, регламентирующих использование ИИ в учебной и исследовательской работе	3,1	6,9	15,7	29,6	44,7	1,06	2
2.8. Я уже изменил(-а) хотя бы один метод обучения или контроля из-за распространения ИИ	5,0	10,1	33,3	33,2	18,4	0,55	0
3. Вызов: Изменение роли преподавателя в триаде «преподаватель – ИИ – обучающийся»							
3.1. Я осознаю, что интеграция ИИ не снижает нагрузку, добавляет педагогу новые функции и ответственность, может способствовать профессиональному выгоранию	4,4	8,2	15,1	31,4	40,9	0,96	2
3.2. Я владею нормативной правовой базой применения ИИ в образовании	16,9	35,1	24,0	17,0	7,0	-0,34	-1
3.3. Я способен(-на) обучить студентов правилам академической этики при использовании ИИ	1,8	5,7	24,5	44,6	23,4	0,78	1
3.4. Я готов(-а) научить студентов использовать профессиональные инструменты ИИ в учебной работе для решения профессиональных задач	11,9	25,0	21,9	22,1	19,1	0,15	-1
3.5. Я готов(-а) организовывать смешанное обучение, сочетая самостоятельную работу студентов с ИИ и аудиторные занятия	12,5	16,9	35,1	23,4	12,1	0	0
3.6. Я готов(-а) обучить студентов выстраивать траектории персонализированного обучения на основе инструментов ИИ	11,8	17,5	36,9	21,1	12,7	-0,05	0
4. Риск: Развитие клипового мышления и снижение когнитивной активности студентов							
4.1. Я замечаю, что студенты всё чаще предпочитают быстрые и готовые ответы от ИИ самостоятельному поиску и анализу материала	1,9	5,0	11,3	29,6	52,2	1,25	2
4.2. Я считаю, что регулярное использование ИИ студентами снижает их способность к самостоятельной постановке проблем и гипотез	3,7	8,2	16,4	31,4	40,3	0,97	2
4.3. Я использую методы (например, критический анализ), препятствующие пассивному копированию студентами ответов генеративного ИИ	11,3	13,8	22,0	34,0	18,9	0,36	1
4.4. Я включаю в учебный процесс задания, при выполнении которых студенты сравнивают материалы обратной связи от разных инструментов ИИ	15,0	25,1	42,0	13,3	4,6	-0,33	0
4.5. Я требую от студентов аргументировать, почему они согласны или не согласны с ответами ИИ	12,5	16,9	25,1	33,2	12,1	0,17	1
4.6. Я модифицирую задания так, чтобы они требовали от студента развёрнутого мыслительного процесса, а не кратного запроса к ИИ	11,9	25,7	34,5	23,9	4,0	-0,18	0
5. Риск: Распространение ИИ-плагиата в академической среде							
5.1. Я знаком(-а) с понятием «ИИ-плагиат» и умею выявлять его признаки	3,8	8,8	18,2	30,2	39,0	0,91	2
5.2. Я объясняю студентам, что присвоение любого материала, сгенерированного ИИ, является нарушением академической этики	1,9	5,1	13,2	31,4	48,4	1,18	2

Продолжение таблицы 2

Утверждение	Варианты ответа (%)					Статистические характеристики	
	-2	-1	0	1	2	Среднее $\bar{x}$	Мода (M_o)
5.3. В моём вузе (на кафедре) есть локальные акты, регулирующие использование студентами ИИ в учебных и исследовательских целях и меры ответственности за нарушение академической этики	28,2	31,4	25,2	7,6	7,6	-0,65	-1
5.4. Я обсуждаю со студентами разницу между помощью ИИ (редактирование текста, поиск идей) и плагиатом	21,9	25,7	24,4	22,7	5,3	-0,08	-2
5.5. Я пересматриваю виды работы и задания, которые наиболее подвержены ИИ-плагиату	12,5	18,2	27,0	31,4	10,9	0,01	1
6. Риск: «Галлюцинации» ИИ, предвзятость и идеологическая ангажированность							
6.1. Я знаю, что ИИ может генерировать ложные и вымышленные сведения («галлюцинировать»)	1,9	4,4	8,8	30,2	54,7	1,29	2
6.2. Я учу студентов перепроверять информацию, полученную от ИИ, по надёжным источникам	11,9	14,4	30,1	32,1	11,5	0,16	0
6.3. Я объясняю, что языковые модели отражают идеологию стран-разработчиков и могут нести предвзятую и идеологически ангажированную информацию	5,0	18,8	27,0	30,8	18,4	0,36	1
6.4. Я требую от студентов использовать разные инструменты ИИ и сравнивать полученную информацию при выполнении учебных заданий и в научной работе	6,7	12,3	35,8	29,9	14,3	0,43	0
6.5. Я требую от студентов подтверждать ключевые фактические утверждения, полученные от ИИ, по авторитетным источникам (учебники, научные статьи и т. п.)	1,9	15,1	41,9	31,4	9,7	0,27	0
7. Риск: Усиление цифрового и социального неравенства							
7.1. Я считаю, что платный доступ к полнофункциональным версиям ИИ усиливает неравенство между студентами	3,8	6,9	33,2	30,2	25,9	0,67	0
7.2. Я предлагаю студентам использовать альтернативные бесплатные инструменты ИИ	11,9	25,0	21,9	34,0	7,2	-0,02	1
7.3. Я готов(-а) адаптировать задания с учётом разного доступа студентов к платным ИИ-инструментам	12,5	16,3	24,5	34,6	12,1	0,17	1
7.4. Я знаю, что в разных вузах принимаются разные решения относительно использования ИИ в обучении и исследовательской работе	2,5	6,9	15,1	32,1	43,4	1,07	2
7.5. Я готов(-а) участвовать в разработке регламентов по использованию ИИ в своём вузе	3,1	16,9	35,7	30,2	14,1	0,33	0
7.6. Я считаю, что разные политики вузов в отношении легализации использования ИИ в учебной деятельности ведут к разделению студентов на тех, кого будут целенаправленно обучать использовать ИИ, и тех, кто будет продолжать несанкционированное заимствование материалов обратной связи от ИИ	5,0	8,8	38,2	29,6	18,4	0,48	0
8. Дополнительные вопросы							
8.1. В целом я чувствую себя готовым(-ой) к системным изменениям в образовании из-за развития ИИ	1,9	5,7	44,4	34,0	14,0	0,54	0
8.2. Я уже адаптирую и модифицирую свои подходы и методы обучения и контроля в ответ на распространение ИИ	2,5	16,9	44,5	34,0	14,0	0,29	0
8.3. Я считаю, что университет должен предоставить мне дополнительное время и ресурсы для работы с ИИ	2,5	6,9	13,2	30,2	47,2	1,13	2

Результаты анкетирования показывают значительный разброс в ответах преподавателей, что свидетельствует об отсутствии на современном этапе в педагогическом сообществе системного понимания основных вызовов и рисков, стоящих перед системой высшего образования, а также готовности педагогов противостоять этим вызовам.

Наиболее высокий уровень согласия респонденты продемонстрировали относительно понимания того, что интеграция ИИ неизбежно приведёт к трансформации профессий и они готовы адаптировать содержание дисциплин под новые профессиональные задачи (В.1.1: $\bar{x} = 1,27$; $M_o = 2$). Около 85% педагогов признают необходимость включения компетенции в области ИИ в число УК студентов всех направлений подготовки (В.1.3: $\bar{x} = 1,05$; $M_o = 2$). Вместе с тем готовность пересмотреть содержание учебных дисциплин со смещением акцента с передачи знаний на развитие практических умений применять знания на практике выражена значительно слабее (В.1.2: $\bar{x} = 0,34$; $M_o = 1$). При этом каждый четвёртый педагог (26,3%), очевидно по причинам потенциальной трудоёмкости, выразил несогласие с данным положением. Наибольшее «сопротивление» со стороны преподавателей вызвал тезис о необходимости регулярного пересмотра и обновления ОПОП из-за появления новых профессиональных инструментов на базе ИИ (В.1.5: $\bar{x} = 0,32$; $M_o = 0$).

В целом преподаватели поддерживают необходимость разработки локальных нормативных документов, регламентирующих использование ИИ в учебном процессе и исследовательской работе студентов (В.2.7: $\bar{x} = 1,06$; $M_o = 2$). Около 60% респондентов выражают согласие с тезисом, что обучающиеся должны предоставлять «цифровой след», подтверждающий внеаудиторное взаимодействие с ИИ (В.2.3: $\bar{x} = 0,59$; $M_o = 2$). Однако фактическая готовность к организационным изменениям со стороны педагогов оказывается значительно ниже. Только 51%

опрошенных организуют внеаудиторную практику студентов с ИИ и интегрировать её в традиционную методику обучения (В.2.2: $\bar{x} = 0,35$; $M_o = 1$). Более того, среднее значение по утверждению «я пересматриваю традиционные формы контроля в связи с использованием студентами ИИ» оказалось отрицательным, и 41,4% выразили несогласие (В.2.4: $\bar{x} = -0,07$). 51,6% педагогов указали, что уже изменили хотя бы один из методов обучения и контроля из-за распространения ИИ, однако достаточно низкие данные средней величины и моды свидетельствуют о немассовом и несистемном характере данных изменений (В.2.8: $\bar{x} = 0,55$; $M_o = 0$). Во многом такая позиция безразличия или отрицания связана, с одной стороны, с некомпетентностью преподавателей в методике обучения профильным дисциплинам на основе средств генеративного ИИ или других инструментов ИИ, а с другой – потенциальной трудоёмкостью перехода с традиционных на инновационные методы обучения.

Подавляющее большинство респондентов осознаёт, что интеграция ИИ не снижает, а, наоборот, добавляет педагогу новые функции и ответственность, что может способствовать профессиональному выгоранию (В.3.1: $\bar{x} = 0,96$; $M_o = 2$). При этом отношение к разным новым функциям у педагогов тоже разное. Лишь 24% преподавателей высшей школы владеют нормативной базой применения ИИ в образовании (В.3.2: $\bar{x} = -0,34$; $M_o = -1$). Готовность обучать студентов правилам соблюдения принципов академической этики при использовании средств генеративного ИИ выражена достаточно высоко (В.3.3: $\bar{x} = 0,78$; $M_o = 1$), тогда как способность научить студентов использовать профессиональные инструменты ИИ оценивается значительно ниже (В.3.4: $\bar{x} = 0,15$; $M_o = -1$). Средние значения по утверждениям относительно организации смешанного формата обучения и выстраивания тра-

екторий персонализированного обучения приближаются к нулю (В.3.5: $\bar{x} = 0$; В.3.6: $\bar{x} = -0,05$). Такие данные свидетельствуют о неопределённости педагогов относительно реализации данных функций и об отрицательном отношении к ним по причине трудо-затратности.

Большая часть преподавателей (81,8%) признают, что обучающиеся всё чаще предпочитают быстрые готовые ответы от ИИ самостоятельному поиску, познанию и анализу учебного материала (В.1.4: $\bar{x} = 1,25$; $M_o = 2$). Более 71% опрошенных согласны с тезисом, что регулярное использование ИИ снижает способность студентов к самостоятельной постановке проблем и гипотез (В.4.2: $\bar{x} = 0,97$; $M_o = 2$). Однако практические меры противодействия данному риску применяются недостаточно часто и последовательно. Лишь 18% педагогов включают в учебный процесс задания на сравнение материалов обратной связи от разных инструментов ИИ (В.4.4: $\bar{x} = -0,33$; $M_o = 0$), и 28% готовы модифицировать задания, требующие развёрнутого мыслительного процесса (В.4.6: $\bar{x} = -0,18$; $M_o = 0$).

В целом преподаватели демонстрируют высокий уровень информированности о феномене ИИ-плагиата (В.5.1: $\bar{x} = 0,91$; $M_o = 2$) и говорят студентам, что присвоение сгенерированного ИИ материала является нарушением академической этики (В.5.2: $\bar{x} = 1,18$; $M_o = 2$). В то же время лишь 15,2% респондентов указали на наличие в их вузах локальных актов, регулирующих сферу и степень использования студентами ИИ (В.5.3: $\bar{x} = -0,65$; $M_o = -1$). Более того, среднее значение по утверждению «я обсуждаю со студентами разницу между помощью ИИ и плагиатом» также оказалось отрицательным (В.5.4: $\bar{x} = -0,08$; $M_o = -2$), а пересмотр видов работы, наиболее подверженных ИИ-плагиату, носит эпизодический характер (В.5.5: $\bar{x} = 0,01$).

Почти 85% преподавателей осведомлены о способности ИИ к «галлюцинациям» (В.6.1: $\bar{x} = 1,29$; $M_o = 2$). Однако лишь половина из них обучает студентов перепроверять полученную от ИИ информацию по надёжным источникам. Рекомендации по использованию разных инструментов ИИ для сравнения информации дают 44,2% преподавателей (В.6.4: $\bar{x} = 0,43$; $M_o = 0$). Наибольший разрыв между осознанием риска и готовностью к действию фиксируется в В.6.5 – 41,1% преподавателей последовательно требует подтверждения студентами ключевых фактических утверждений, полученных от ИИ, по авторитетным источникам.

Более половины респондентов (56,1%) согласны, что платный доступ к полнофункциональным версиям ИИ усиливает неравенство между пользователями (В.7.1: $\bar{x} = 0,67$; $M_o = 0$). При этом 75,5% педагогов предполагает, что в разных вузах действуют разные политики относительно использования ИИ в обучении и исследовательской работе (В.7.4: $\bar{x} = 1,07$; $M_o = 2$), и 48% считают, что такая дифференциация ведёт к разделению студентов (В.7.6: $\bar{x} = 0,48$; $M_o = 0$). Готовность адаптировать задания с учётом разного доступа студентов к платным инструментам выразили 46,7% опрошенных (В.7.3: $\bar{x} = 0,17$; $M_o = 1$), а участвовать в разработке регламентов по использованию ИИ в своём вузе готовы 44,3% опрошенных (В.7.5: $\bar{x} = 0,33$; $M_o = 0$).

Отвечая на заключительные дополнительные вопросы, лишь 48% преподавателей выразили готовность к системным изменениям в образовании из-за развития ИИ (В.8.1: $\bar{x} = 0,54$; $M_o = 0$). Ещё меньше (34%) уже адаптируют и модифицируют свои подходы и методы обучения и контроля в условиях распространения ИИ (В.8.2: $\bar{x} = 0,29$; $M_o = 0$). При этом 77,4% респондентов считают, что университет должен предоставить им дополнительное время и ресурсы для работы с ИИ (В.8.3: $\bar{x} = 1,13$; $M_o = 2$).

Обсуждение результатов

Полученные в ходе онлайн-анкетирования данные позволили *выявить* не только текущее понимание преподавателями высшей школы вызовов и рисков, связанных с интеграцией ИИ в образование, но и, главное, *ряд системных противоречий*, требующих осмысления и ответных действий со стороны руководителей систем образования как локального, так и федерального уровня. Результаты исследования фиксируют ситуацию «адаптивного консерватизма», при котором на уровне понимания и декларирования признаётся значимость изменений, а на уровне практической деятельности наблюдается высокая доля инертности. Выделим шесть аспектов для научной дискуссии, которые будут определять перспективы дальнейших исследований.

Во-первых, материалы работы свидетельствуют о наличии *системного разрыва между пониманием педагогами* основных вызовов и рисков, стоящих перед высшей школой в эпоху распространения генеративного ИИ, *и их практической деятельностью*, направленной на сокращение возникших угроз. Подтверждением могут служить ответы на вопросы каждого из разделов анкеты (например, В.4.1, 4.2 vs В.4.4 и 4.6; В. 6.1 vs В.6.2; В.7.4, 7.6 vs В.7.3). В среднем процент педагогов, активно использующих ИИ в обучении профильным дисциплинам на момент проведения исследования, составляет 25–30%. Эти результаты соотносятся с данными, полученными в 2023 г. П.В. Сысоевым [36] в ходе эмпирического исследования по выявлению доли преподавателей вузов, интегрирующих инструменты ИИ в свою работу. Подобные сходства в количественных данных опросов оказались несколько неожиданными. За три года распространения ИИ не наблюдается резкого скачка в интеграции преподавателями высшей школы ИИ инструментов в свою профессиональную деятельность. Без наличия вертикального управления со стороны Минобрнауки системного внедрения ИИ в образование не произошло. Педагоги-инно-

ваторы, которые с самого начала появления *ChatGPT* в 2022 г. начали изучать его дидактический потенциал, продолжают применять разные инструменты ИИ для решения учебных, методических и исследовательских задач. Их коллеги, которые отрицали ИИ или предпочитали не использовать инструменты ИИ, продолжают это делать, тем самым не способствуя противостоянию вызовам и рискам. С одной стороны, этот разрыв может быть объяснён общей некомпетентностью в сфере ИИ в образовании большого числа преподавателей и их нежеланием брать на себя дополнительные трудоёмкие и затратные по времени функции. С другой – это может быть *формированием специфической институциональной защиты от постоянных инноваций*. Обсуждение со студентами принципов академической этики без значительных изменений в используемых подходах и методах обучения или осознание важности компетенции обучающихся в области ИИ и её отсутствие в большинстве ОПОП приводят к поверхностным и формальным изменениям в образовании (без трансформации ядра профессиональной деятельности), неспособным качественно ответить на стоящие вызовы и риски. Полученные данные перекликаются с выводами аналитических исследований западных коллег. В частности, Р. Тахери, Н. Дажеми и соавторы [37] подчёркивают, что реакция системы образования на ИИ часто носит характер «адаптивного консерватизма», когда институциональные механизмы изменяются медленнее, чем когнитивное принятие новой реальности преподавателями. Результаты проведённого исследования являются эмпирическим подтверждением данного тезиса.

Во-вторых, по многим аспектам интеграции ИИ в образование наблюдается *консервация «зон неопределённости» как некий ответ на постоянные инновации*. Обращает на себя внимание высокая доля нейтральных ответов по утверждениям, связанным с пересмотром содержания обучения, подходов и методов обучения, разработкой локальных

нормативных документов и внедрением форм смешанного обучения. Нейтральная позиция по большому числу вопросов может свидетельствовать об определённой зоне автономии преподавателей от административного давления. Отсутствие федеральных нормативных документов и локальных актов, регламентирующих легальное использование педагогами и студентами средств ИИ, не только не стимулирует изменения, а, наоборот, легитимирует бездействие преподавателей и администраторов среднего звена (заведующих кафедрами, руководителей ОПОП, деканов, начальников учебных управлений). Значительный разрыв между знаниями и практической деятельностью по применению ИИ в учебном процессе через свои ответы педагоги объясняют отсутствием «официальных правил игры». Данная ситуация, на наш взгляд, обостряет многие риски. Чем дольше система откладывает нормативное оформление использования ИИ, тем более изощрёнными становятся неформальные практики скрытого использования ИИ студентами, процветание ИИ-плагиата и дальнейшее социальное неравенство обучающихся.

В-третьих, результаты проведённого исследования подтверждают гипотезу, выдвинутую Я.И. Кузьминовым, Е.В. Кручинской, И.А. Груздевым и А.А. Наумовым [20], о том, что *основное социальное разделение в эпоху ИИ будет определяться* не столько возможностью обучающихся оплатить доступ к платным ресурсам на базе ИИ, сколько *институциональной культурой работы с технологическими решениями на основе ИИ*. 75,5% (В.7.4) преподавателей согласны, что в разных вузах действуют разные политики в отношении ИИ, и 48% видят в этом фактор дифференциации студентов. Однако, на наш взгляд, наиболее тревожным представляется другое. Как показывает качественный анализ материалов анкетирования, даже в рамках одного вуза разброс в ответах преподавателей по некоторым вопросам огромен (от полного несогласия до активного вне-

дрения). Это означает, что социальное неравенство происходит уже на уровне кафедр. Студенты одного направления подготовки могут получать принципиально разный опыт работы с ИИ в зависимости от того, какой преподаватель ведёт дисциплину. Возможным перспективным направлением для дальнейшего изучения может быть рассмотрение того, как именно происходит эта «микросегрегация» в рамках кафедры, по каким критериям преподаватели разделяются на «инноваторов» и «консерваторов», по каким критериям «инноваторы» отбирают те или иные инструменты ИИ, как они их используют в учебном процессе и как всё это влияет на развитие учебной автономии студентов.

В-четвёртых, соотнесение педагогического дискурса об интеграции ИИ в образование с результатами анкетирования показывает *парадокс «педагогического восторга», переходящего в профессиональное выгорание*. В ответах респондентов чётко фиксируется противоречивое отношение к новым функциям, которые приобретает преподаватель в эпоху распространения ИИ. С одной стороны, 85% респондентов осознают, что интеграция ИИ добавляет педагогам новые функции и ответственность (В.3.1). С другой – готовность осваивать конкретные новые функции (организация смешанного обучения, выстраивание траекторий персонализированного обучения, использование профессиональных инструментов ИИ и т. п.) оценивается значительно ниже, вплоть до отрицательных средних значений (В.3.6: $\bar{x} = -0,05$). Это свидетельствует об определённом когнитивном диссонансе: преподаватель признаёт неизбежность изменений, но или не видит, или не хочет видеть способов их реализации без существенных изменений в системе нормирования труда, оплаты и повышения квалификации. Более того, сама риторика о «дидактическом потенциале ИИ», преобладающая в научной литературе, может парадоксальным образом усиливать выгорание педагогов. Она создаёт

завышенные ожидания относительно лёгкости интеграции ИИ и исключительных достоинств инструментов ИИ в решении учебных или исследовательских задач, в то время как реальность будет требовать значительных временных и когнитивных затрат со стороны преподавателя. Данная позиция коррелирует с выводами Д. Коттона, П. Коттона и Дж. Шипвея [16] о том, что внедрение ИИ парадоксальным образом увеличивает, а не снижает когнитивную нагрузку на педагога, который, наряду с традиционными функциями, должен выступать «медиатором ИИ» и «верификатором информации».

В-пятых, особого внимания заслуживает обсуждение вопросов, связанных с *распространением в академической среде ИИ-плагиата*. Как свидетельствуют результаты анкетирования, налицо разрыв между декларируемой позицией преподавателей по этому вопросу и реальными практиками. Несмотря на то, что 79,8% респондентов говорят студентам о недопустимости присвоения авторства сгенерированных ИИ текстов (В.5.2), только 28% обсуждают со студентами различия между использованием генеративного ИИ в качестве помощника и ИИ-плагиатом, и лишь 15,2% работают в вузах, где есть локальные акты, регламентирующие использование ИИ в учебной и исследовательской работе. Эта ситуация создаёт институциональную ловушку. С одной стороны, обучающиеся получают посыл от преподавателей, что ИИ-плагиат – это плохо. С другой – в большинстве вузов не существует инструментов ответственности за неэтичное использование ИИ. В результате, как показывают данные эмпирического исследования П.В. Сысоева [17], около половины студентов начинают воспринимать заимствование материалов обратной связи от ИИ в качестве нормы. Перспектива исследований в данном направлении видится в дальнейшем изучении того, как именно формируются договорённости между преподавателями и студентами по легальному использованию ИИ, как изменение подхо-

дов, методов и форм преподавания и проведения контроля влияет на распространение ИИ-плагиата.

В-шестых, перспектива системного использования ИИ в образовании определяется системной политикой. Полученные в ходе настоящего исследования данные, а также результаты работ других учёных в области интеграции ИИ в образование [38–40] ставят вопрос о необходимости изменения самого подхода – *от точечной готовности отдельных преподавателей использовать конкретные инструменты ИИ в своей работе к системной институциональной политике по комплексному использованию ИИ в образовании*. Достаточно низкий процент педагогов, использующих ИИ на регулярной основе, и практически отсутствие изменений в количестве по сравнению с результатами исследования 2023 г. [36] говорят об *исключительно инициативном использовании преподавателями ИИ в профессиональной работе*. Отметим, что в целом система образования является негибким и консервативным институтом, позволяющим изменения, проводимые исключительно вертикально сверху (например, переход на двухуровневую модель подготовки специалистов или переход на дистанционную форму обучения в период пандемии и т. п.). Отсутствие федеральных нормативных документов, регулирующих сферу и степень использования студентами средств генеративного ИИ в учебной и исследовательской работе, обуславливает позицию выжидания со стороны руководства большинства вузов страны. Отсюда факультативное и несистемное использование потенциала ИИ педагогами в ходе преподавания профильных дисциплин. Исключение составляют вузы, в которых приняты локальные нормативные акты и разрабатываются методические системы комплексной подготовки специалистов на основе интеграции ИИ [21; 32; 34; 35]. Очевидна необходимость введения в ФГОС ВО универсальной компетенции в области ИИ, а также разработки Минобрнауки типово-

го регламента использования ИИ в учебных и исследовательских целях (по аналогии с антиплагиатом). Перспектива дальнейших исследований видится в изучении того, как различные институциональные инструменты (финансовые стимулы, административные санкции и т. п.) влияют на преодоление описанного разрыва и какие ответы на вызовы ИИ предлагают разные типы вузов (национальные исследовательские *vs* региональные).

Ограничения исследования

Интерпретация данных настоящего исследования должна проводиться с учётом ряда ограничений. Во-первых, несмотря на репрезентативную выборку и комплексный дизайн инструментария, все данные основаны на самооценке респондентов. Это могло создать риск социально желательных ответов, особенно в вопросах, касающихся готовности и способности преподавателей противостоять вызовам и рискам.

Во-вторых, в исследовании приняли участие преподаватели вузов, разных по своей миссии и целевой направленности (национальных исследовательских университетов и региональных вузов). Объединённая статистика показывает усреднённое значение и нивелирует различия между показателями разных типов вузов. Отдельное исследование может быть направлено на изучение сходств и различий в способности интегрировать ИИ в образовательную деятельность представителями разных типов вузов.

В-третьих, в анкетировании участвовали преподаватели разных профильных дисциплин, работающие на разных направлениях подготовки или специальностях. По объективным причинам сфера и степень использования технологических решений на базе ИИ в зависимости от направления подготовки могут изменяться.

Данные ограничения определяют область дальнейших перспективных исследований в области интеграции ИИ в образование.

Заключение

Результаты исследования показали наличие системного разрыва между декларируемым пониманием вызовов и рисков и реальной практической готовностью к организационным и методическим изменениям. Порядка 30–40% респондентов выразили нейтральное отношение по большинству вопросов, связанных с внедрением новых методов и форм контроля, а 15–25% опрошенных – отрицательное отношение. Это свидетельствует о том, что системная интеграция ИИ в высшее образование не может осуществляться преимущественно на инициативной основе отдельными преподавателями. Противостояние вызовам и рискам ИИ возможно исключительно на институциональном уровне при разработке и внедрении нормативных регламентов использования ИИ, введении универсальной компетенции в области ИИ во ФГОС ВО, пересмотре учебной нагрузки преподавателей и создании ресурсной базы для персонализированной работы студентов с инструментами ИИ.

Данная работа может положить начало ряду дальнейших исследований, направленных на: а) изучение механизмов сопротивления, позволяющих педагогам обосновывать своё бездействие; б) сравнение вузов различного типа (национальные исследовательские *vs* региональные), а также проведение лонгитюдных исследований изменений отношения преподавателей к интеграции ИИ и их практике в рамках конкретных вузов; в) разработку и апробацию организационных моделей интеграции ИИ в систему высшего образования в конкретных вузах.

Литература

1. Сысоев П.В., Филатов Е.М., Сорокин Д.О. Обратная связь в обучении иностранному языку: от информационных технологий к искусственному интеллекту // Язык и культура. – 2024. – № 65. – С. 242–261. – DOI: 10.17223/19996195/65/11.
2. Chan K., Zary N. Applications and Challenges of Implementing Artificial Intelligence in Medical Education: Integrative Review // JMIR Medical

- Education. – 2019. – Vol. 5, no. 1. – Article no. 13930. – DOI: 10.2196/13930.
3. Итинсон К.С. Информатизация медицинского образования: системы искусственного интеллекта в обучении студентов и врачей // Балтийский гуманитарный журнал. – 2020. – Т. 9, № 3 (32). – С. 91–93. – DOI: 10.26140/bgз3-2020-0903-0021.
 4. Zhang W., Cai M., Lee H., Evans R., Zhu C., Ming C. AI in Medical Education: Global situation, effects and challenges // Education and Information Technologies. – 2024. – Vol. 29. – P. 4611–4633. – DOI: 10.1007/s10639-023-12009-8.
 5. Ягудина А.Р., Цилицкий В.С., Виноградова И.В., Кузнецова С.Б., Жарина Н.А. Искусственный интеллект и его роль в преподавании экономических дисциплин в вузе // Московский экономический журнал. – 2022. – № 2. – С. 634–642. – DOI: 10.55186/2413046X_2022_7_2_104.
 6. Feuerriegel S., Shrestha Y.R., von Krogh G., Zhang C. Bringing artificial intelligence to business management // Nature Machine Intelligence. – 2022. – Vol. 4, no. 7. – P. 611–613. – DOI: 10.1038/s42256-022-00512-5.
 7. Иванова А.П. Искусственный интеллект в сфере права и юридической практике: Основные проблемы и перспективы развития // Социальные и гуманитарные науки. Отечественная и зарубежная литература. Серия 4: Государство и право. – 2021. – № 1. – С. 90–98. – DOI: 10.31249/rgpravo/2021.01.09.
 8. Сысоев П.В., Гаврилов М.В., Булочников С.Ю. Матрица технических решений на базе искусственного интеллекта в профессиональной подготовке будущих юристов // Вестник Тамбовского университета. Серия: Гуманитарные науки. – 2025. – Т. 30, № 2. – С. 336–351. – DOI: 10.20310/1810-0201-2025-30-2-336-351.
 9. Смык А.Ф., Гусева Е.А. Искусственный интеллект в образовании будущих инженеров: фокус на предметной области // Инженерное образование. – 2025. – № 38. – С. 205–221. – DOI: 10.54835/18102883_2025_38_18.
 10. Ayeni O.O., Hamad N.M.A., Chisom O.N., Osawaru B., Adewusi O.E. AI in education: A review of personalized learning and educational technology // GSC Advanced Research and Reviews. – 2024. – No. 18 (02). – P. 261–271. – DOI: 10.30574/gscarr.2024.18.2.0062.
 11. Jegede O.O. Artificial Intelligence and English Language Learning: Exploring the Roles of AI-Driven Tools in Personalizing Learning and Providing Instant Feedback // Universal Library of Languages and Literatures. – 2024. – No. 1 (2). – P. 6–19. – DOI: 10.70315/uloap.ulli.2024.0102002.
 12. Сысоев П.В. Персонализированное обучение на основе технологий искусственного интеллекта: насколько готовы современные студенты к новым возможностям получения образования // Высшее образование в России. – 2025. – Т. 34, № 2. – С. 51–71. – DOI: 10.31992/0869-3617-2025-34-2-51-71.
 13. Титова С.В. Интеллектуальные системы обучения для персонализации и адаптации языковых курсов // Вестн. Моск. ун-та. Сер. 19. Лингвистика и межкультурная коммуникация. – 2024. – Т. 27, № 4. – С. 84–99. – DOI: 10.55959/MSU-2074-1588-19-27-4-6.
 14. Казакова Е.И., Кузьминов Я.И. Мы должны воспитать культуру критического отношения к ответам искусственного интеллекта // Вопросы образования. – 2025. – № 1. – С. 8–24. – DOI: 10.17323/vo-2025-25882.
 15. Ивахненко Е.Н., Никольский В.С. ChatGPT в высшем образовании и науке: угроза или ценный ресурс? // Высшее образование в России. – 2023. – Т. 32, № 4. – С. 9–22. – DOI: 10.31992/0869-3617-2023-32-4-9-22.
 16. Cotton D.R.E., Cotton P.A., Shipway J.R. Chatting and cheating: Ensuring academic integrity in the era of ChatGPT // Innovations in Education and Teaching International. – 2023. – DOI: 10.1080/14703297.2023.2190148.
 17. Сысоев П.В. Этика и ИИ-плагиат в академической среде: понимание студентами вопросов соблюдения авторской этики и проблемы плагиата в процессе взаимодействия с генеративным искусственным интеллектом // Высшее образование в России. – 2024. – Т. 33, № 2. – С. 31–53. – DOI: 10.31992/0869-3617-2024-33-2-31-53.
 18. Титова С.В., Харламенко И.В. Структура профессиональной компетенции педагога иностранных языков в области использования искусственного интеллекта // Язык и культура. – 2025. – № 69. – С. 220–246. – DOI: 10.17223/19996195/69/11.
 19. Сысоев П.В. Компетенция современного педагога в области искусственного интеллекта: структура и содержание // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 58–79. – DOI: 10.31992/0869-3617-2025-34-6-58-79.

20. Кузьминов Я.И., Кручинская Е.В., Груздев И.А., Наумов А.А. Отстающие и опережающие: как студенты используют генеративный искусственный интеллект в образовательных целях // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 9–35. – DOI: 10.31992/0869-3617-2025-34-6-9-35.
21. Сысоев П.В., Евстигнеев М.Н. Интеграция технологий искусственного интеллекта в лингвометодическую подготовку будущих учителей иностранного языка // Язык и культура. – 2025. – № 69. – С. 204–219. – DOI: 10.17223/19996195/69/10.
22. Сысоев П.В., Филатов Е.М. Методика обучения студентов написанию иноязычных творческих работ на основе оценочной обратной связи от искусственного интеллекта // Перспективы науки и образования. – 2024. – № 1 (67). – С. 115–135. – DOI: 10.32744/pse.2024.1.6.
23. Сорокин Д.О. Развитие умений устного иноязычного взаимодействия на основе практики обучающихся с инструментами искусственного интеллекта // Иностранные языки в школе. – 2026. – № 2. – С. 31–41. – EDN: JLIPIHY.
24. Филатов Е.М. Развитие умений написания структурных компонентов эссе на основе оценочной и корректирующей обратной связи от искусственного интеллекта // Иностранные языки в школе. – 2026. – № 2. – С. 42–52. – EDN: JMKENC.
25. Никольский В.С. Коммуникативный искусственный интеллект: концептуализация новой реальности в образовании // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 152–168. – DOI: 10.31992/0869-3617-2025-34-6-152-168.
26. Waltzer T., Dahl A. Why do students cheat? Perceptions, evaluations, and motivations // Ethics & Behavior. – 2023. – Vol. 33, no. 2. – P. 130–150. – DOI: 10.1080/10508422.2022.2026775.
27. Сысоев П.В., Филатов Е.М. ChatGPT в исследовательской работе студентов: запрещать или обучать? // Вестник Тамбовского университета. Серия: Гуманитарные науки. – 2023. – Т. 28, № 2. – С. 276–301. – DOI: 10.20310/1810-0201-2023-28-2-276-301.
28. Kohnke L., Moorhouse B.L., Zou D. ChatGPT for language teaching and learning // RELC Journal. – 2023. – Vol. 54, no. 2. – P. 537–550. – DOI: 10.1177/00336882231162868.
29. Dwivedi Y.K., Kshetri N., Hughes L., Slade E.L., Jeyaraj A. et al. Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy // International Journal of Information Management. – 2023. – Vol. 71. – Article no. 102642. – DOI: 10.1016/j.ijinfomgt.2023.102642.
30. Kasneci E., Sessler K., Küchemann S., Bannert M., Dementieva D. et al. ChatGPT for good? On opportunities and challenges of large language models for education // Learning and Individual Differences. – 2023. – Vol. 103. – Article no. 102274. – DOI: 10.1016/J.LINDIF.2023.102274.
31. Bernabei M., Colabianchi S., Falegnami A., Costantino F. Students’ use of large language models in engineering education: A case study on technology acceptance, perceptions, efficacy, and detection chances // Computers and Education: Artificial Intelligence. – 2023. – Vol. 5. – Article no. 100172. – DOI: 10.1016/j.caeai.2023.100172.
32. Тивьяева И.В., Михайлова С.В., Казанцева А.А. Регламентирование использования средств генеративного искусственного интеллекта в выпускной квалификационной работе // Вестник МГПУ. Серия «Филология. Теория языка. Языковое образование». – 2024. – № 2 (54). – С. 202–218. – DOI: 10.25688/2076-913X.2024.54.2.15.
33. Сысоев П. В., Евстигнеев М. Н. Использование студентами педагогических специальностей технических решений на основе искусственного интеллекта в ходе педагогической практики // Перспективы науки и образования. – 2025. – № 4. – С. 135–150. – DOI: 10.32744/pse.2025.4.9.
34. Ананин Д.П., Комаров Р.В., Реморенко И.М. «Когда честно – хорошо, для имитации – плохо»: стратегии использования генеративного искусственного интеллекта в российском вузе // Высшее образование в России. – 2025. – Т. 34, № 2. – С. 31–50. – DOI: 10.31992/0869-3617-2025-34-2-31-50.
35. Сысоев П.В. Экосистема подготовки будущих учителей на основе технологий искусственного интеллекта // Перспективы науки и образования. – 2026. – № 1. – С. 642–658. – DOI: 10.32744/pse.2026.1.40
36. Сысоев П.В. Искусственный интеллект в образовании: осведомлённость, готовность и практика применения преподавателями выс-

- шей школы технологий искусственного интеллекта в профессиональной деятельности // Высшее образование в России. – 2023. – Т. 32, № 10. – С. 9–33. – DOI: 10.31992/0869-3617-2023-32-10-9-33.
37. Taheri R., Nazemi N., Pennington S.E., Clark J.A., Dadgostari F. Factors influencing educators' ai adoption: A grounded meta-analysis review // Computers and Education: Artificial Intelligence. – 2025. – No. 9. – Article no. 100464. – DOI: 10.1016/j.caeai.2025.100464.
 38. Rensfeldt A.B., Rahm L. Automating Teacher Work? A History of the Politics of Automation and Artificial Intelligence in Education // Postdigital Science and Education. – 2023. – No. 5. – P. 25–43. – DOI: 10.1007/s42438-022-00344-x.
 39. Gang D., Yufeng S., Zhao Y. The Innovation of Ideological and Political Education Integrating Artificial Intelligence Big Data with the Support of Wireless Network // Applied Artificial Intelligence. – 2023. – Vol. 37, no. 1. – Article no. 2219943. – DOI: 10.1080/08839514.2023.2219943.
 40. Bulathwela S., Pírez-Ortiz M., Holloway C., Cukurova M., Shawe-Taylor J. Artificial Intelligence Alone Will Not Democratise Education: On Educational Inequality, Techno-Solutionism and Inclusive Tools // Sustainability. – 2024. – Vol. 16, no. 2. – Article no. 781. – DOI: 10.3390/su16020781.
- Статья поступила в редакцию 29.04.2026
Принята к публикации 29.05.2026

References

1. Sysoyev, P.V., Filatov, E.M., Sorokin, D.O. (2024). Feedback in Foreign Language Teaching: From Information Technologies to Artificial Intelligence. *Yazyk i kultura = Language and Culture*. Vol. 65, pp. 242–261, doi: 10.17223/19996195/65/11 (In Russ., abstract in Eng.).
2. Chan, K., Zary, N. (2019). Applications and Challenges of Implementing Artificial Intelligence in Medical Education: Integrative Review. *JMIR Medical Education*. Vol. 5, no. 1, article no. 13930, doi: 10.2196/13930.
3. Itinson, K.S. (2020). Informatization of Medical Education: Artificial Intelligence Systems in Training Students and Doctors. *Baltiyskiy gumanitarnyy zhurnal = Baltic Humanitarian Journal*. Vol. 9, no. 3 (32), pp. 91–93, doi: 10.26140/bg3-2020-0903-0021 (In Russ., abstract in Eng.).
4. Zhang, W., Cai, M., Lee, H., Evans, R., Zhu, C., Ming, C. (2024). AI in Medical Education: Global situation, effects and challenges. *Education and Information Technologies*. Vol. 29, pp. 4611–4633, doi: 10.1007/s10639-023-12009-8.
5. Yagudina, A.R., Tsilitzky, V.S., Vinogradova, I.V., Kuznetsova, S.B., Zharina, N.A. (2022). Artificial Intelligence and Its Role in the Teaching of Economic Disciplines at the University. *Moskovskiy ekonomicheskyy zhurnal = Moscow Economic Journal*. No. 2, pp. 634–642, doi: 10.55186/2413046X_2022_7_2_104 (In Russ., abstract in Eng.).
6. Feuerriegel, S., Shrestha, Y.R., von Krogh, G., Zhang, C. (2022). Bringing Artificial Intelligence to Business Management. *Nature Machine Intelligence*. Vol. 4, no. 7, pp. 611–613, doi: 10.1038/s42256-022-00512-5.
7. Ivanova, A.P. (2021). Artificial Intelligence in the Field of Law and Legal Practice: The Main Problems and Prospect of Development. *Sotsialnyye i gumanitarnyye nauki. Otechestvennaya i zarubezhnaya literatura. Seriya 4: Gosudarstvo i pravo = Social Sciences and Humanities. Domestic and Foreign Literature. Series 4: State and Law*. No. 1, pp. 90–98, doi: 10.31249/rg-pravo/2021.01.09 (In Russ., abstract in Eng.).
8. Sysoyev, P.V., Gavrilov, M.V., Bulochnikov, S.Yu. (2025). Matrix of Technical Solutions Based on Artificial Intelligence in the Professional Training of Future Lawyers. *Vestnik Tambovskogo universiteta. Seriya: Gumanitarnyye nauki = Tambov University Review. Series: Humanities*. Vol. 30, no. 2, pp. 336–351, doi: 10.20310/1810-0201-2025-30-2-336-351 (In Russ., abstract in Eng.).

9. Smyk, A.F., Guseva, E.A. (2025). Artificial Intelligence in the Education of Future Engineers: Focus on the Subject Area. *Inzhenernoe obrazovanie = Engineering Education*. No. 38, pp. 205-221, doi: 10.54835/18102883_2025_38_18 (In Russ., abstract in Eng.).
10. Ayeni, O.O., Hamad, N.M.A., Chisom, O.N., Osawaru, B., Adewusi, O.E. (2024). AI in Education: A Review of Personalized Learning and Educational Technology. *GSC Advanced Research and Reviews*. No. 18 (02), pp. 261-271, doi: 10.30574/gscarr.2024.18.2.0062.
11. Jegede, O.O. (2024). Artificial Intelligence and English Language Learning: Exploring the Roles of AI-Driven Tools in Personalizing Learning and Providing Instant Feedback. *Universal Library of Languages and Literatures*. No. 1 (2), pp. 6-19, doi: 10.70315/uloap.ulli.2024.0102002.
12. Sysoyev, P.V. (2025). Personalized Learning Based on Artificial Intelligence: How Ready Are Modern Students for New Educational Opportunities. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 2, pp. 51-71, doi: 10.31992/0869-3617-2025-34-2-51-71 (In Russ., abstract in Eng.).
13. Titova, S.V. (2024). Intelligent Learning Systems for Personalizing and Adapting Language Courses. *Vestnik Moskovskogo universiteta. Seriya 19. Lingvistika i mezhkul'turnaya kommunikatsiya = Moscow University Bulletin. Series 19. Linguistics and Intercultural Communication*. Vol. 27, no. 4, pp. 84-99, doi: 10.55959/MSU-2074-1588-19-27-4-6 (In Russ., abstract in Eng.).
14. Kazakova, E.I., Kuzminov, Ya.I. (2025). We Should Foster a Culture of Critical Attitude toward Artificial Intelligence. *Voprosy obrazovaniya = Educational Studies Moscow*, No. 1, pp. 8-24, doi: 10.17323/vo-2025-25882 (In Russ., abstract in Eng.).
15. Ivakhnenko, E.N., Nikolskiy, V.S. (2023). ChatGPT in Higher Education and Science: A Threat or a Valuable Resource? *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 4, pp. 9-22, doi: 10.31992/0869-3617-2023-32-4-9-22 (In Russ., abstract in Eng.).
16. Cotton, D.R.E., Cotton, P.A., Shipway, J.R. (2023). Chatting and Cheating: Ensuring Academic Integrity in the Era of ChatGPT. *Innovations in Education and Teaching International*, doi: 10.1080/14703297.2023.2190148.
17. Sysoyev, P.V. (2024). Ethics and AI-Plagiarism in an Academic Environment: Students' Understanding of Compliance with Author's Ethics and the Problem of Plagiarism in the Process of Interaction with Generative Artificial Intelligence. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 33, no. 2, pp. 31-53, doi: 10.31992/0869-3617-2024-33-2-31-53 (In Russ., abstract in Eng.).
18. Titova, S.V., Kharlamenko, I.V. (2025). The Framework of Professional Competence for Foreign Language Teachers Utilizing Artificial Intelligence. *Yazyk i Kultura = Language and Culture*. Vol. 69, pp. 220-246, doi: 10.17223/19996195/69/11 (In Russ., abstract in Eng.).
19. Sysoyev, P.V. (2025). A Modern Teacher's Competence in the Field of Artificial Intelligence: Structure and Content. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 58-79, doi: 10.31992/0869-3617-2025-34-6-58-79 (In Russ., abstract in Eng.).
20. Kuzminov, Ya.I., Kruchinskaia, E.V., Gruzdev, I.A., Naumov, A.A. (2025). Falling Behind and Getting Ahead: Student Use of Generative AI in Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 9-35, doi: 10.31992/0869-3617-2025-34-6-9-35 (in Russ., abstract in Eng.).
21. Sysoyev, P.V., Evstigneev, M.N. (2025). Integration of Artificial Intelligence Technologies in Language and Methodological Pre-Service Teachers' Training. *Yazyk i kul'tura = Language and Culture*. No. 69, pp. 204-219, doi: 10.17223/19996195/69/10 (In Russ., abstract in Eng.).
22. Sysoyev, P.V., Filatov, E.M. (2024). Method of Teaching Students' Foreign Language Creative Writing Based on Evaluative Feedback from Artificial Intelligence. *Perspektivy nauki i*

- obrazovania – Perspectives of Science and Education*. Vol. 67 (1), pp. 115-135, doi: 10.32744/pse.2024.1.6 (In Russ., abstract in Eng.).
23. Sorokin, D.O. (2026). Developing Oral Foreign Language Interaction Skills Based on Students' Practice with Artificial Intelligence Tools. *Inostrannye yazyki v shkole = Foreign Languages at School*. No. 2, pp. 31-41. Available at: <https://elibrary.ru/item.asp?id=89013681> (accessed 20.04.2026). (In Russ., abstract in Eng.).
 24. Filatov, E.M. (2026). Developing Skills in Writing Structural Components of Essays Based on Evaluative and Corrective Feedback from Artificial Intelligence. *Inostrannye yazyki v shkole = Foreign Languages at School*. No. 2, pp. 42-52. Available at: <https://elibrary.ru/item.asp?id=89013682> (accessed 20.04.2026). (In Russ., abstract in Eng.).
 25. Nikolskiy, V.S. (2025). Communicative Artificial Intelligence: Conceptualizing a New Reality in Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 152-168, doi: 10.31992/0869-3617-2025-34-6-152-168 (In Russ., abstract in Eng.).
 26. Waltzer, T., Dahl, A. (2023). Why Do Students Cheat? Perceptions, Evaluations, and Motivations. *Ethics & Behavior*. Vol. 33, no. 2, pp. 130-150, doi: 10.1080/10508422.2022.2026775.
 27. Sysoyev, P.V., Filatov, E.M. (2023). ChatGPT in Students' Research Work: To Forbid or to Teach? *Vestnik Tambovskogo universiteta. Seriya: Gumanitarnye nauki = Tambov University Review. Series: Humanities*. Vol. 28, no. 2, pp. 276-301, doi: 10.20310/1810-0201-2023-28-2-276-301 (In Russ., abstract in Eng.).
 28. Kohnke, L., Moorhouse, B.L., Zou, D. (2023). ChatGPT for Language Teaching and Learning. *RELC Journal*. Vol. 54, no. 2, pp. 537-550, doi: 10.1177/00336882231162868.
 29. Dwivedi, Y.K., Kshetri, N., Hughes, L., Slade, E.L., Jeyaraj, A. et al. (2023). Opinion Paper: "So What If ChatGPT Wrote It?" Multidisciplinary Perspectives on Opportunities, Challenges and Implications of Generative Conversational AI for Research, Practice and Policy. *International Journal of Information Management*. Vol. 71, article no. 102642, doi: 10.1016/j.ijinfo-mgt.2023.102642.
 30. Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D. et al. (2023). ChatGPT for Good? On Opportunities and Challenges of Large Language Models for Education. *Learning and Individual Differences*. Vol. 103, article no. 102274, doi: 10.1016/J.LINDIF.2023.102274.
 31. Bernabei, M., Colabianchi, S., Falegnami, A., Costantino, F. (2023). Students' Use of Large Language Models in Engineering Education: A Case Study on Technology Acceptance, Perceptions, Efficacy, and Detection Chances. *Computers and Education: Artificial Intelligence*. Vol. 5, article no. 100172, doi: 10.1016/j.caeai.2023.100172.
 32. Tivyaeva, I.V., Mikhailova, S.V., Kazantseva, A.A. (2024). Regulating the Use of Generative Artificial Intelligence Tools in Graduate Qualification Papers. *MCU Journal of Philology. Theory of Linguistics. Linguistic Education*. No. 2 (54), pp. 202-218, doi: 10.25688/2076-913X.2024.54.2.15.
 33. Sysoyev, P.V., Evstigneev, M.N. (2025). The Use of Technical Solutions Based on Artificial Intelligence in Teaching Internship in Pre-Service Teacher Training Programs. *Perspektivy nauki i obrazovania = Perspectives of Science and Education*. No. 4, pp. 135-150, doi: 10.32744/pse.2025.4.9 (In Russ., abstract in Eng.).
 34. Ananin, D.P., Komarov R.V., Remorenko, I.M. (2025). "When Honesty is Good, for Imitation is Bad": Strategies for Using Generative Artificial Intelligence in Russian Higher Education Institutions. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 2, pp. 31-50, doi: 10.31992/0869-3617-2025-34-2-31-50 (In Russ., abstract in Eng.).

35. Sysoyev, P.V. (2026). An Ecosystem for Teacher Training Based on Artificial Intelligence Technologies. *Perspektivy nauki i obrazovaniya = Perspectives of Science and Education*. No. 1, pp. 642-658, doi: 10.32744/pse.2026.1.40.
36. Sysoyev, P.V. (2023). Artificial Intelligence in Education: Awareness, Readiness and Practice of Using Artificial Intelligence Technologies in Professional Activities by University Faculty. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 10, pp. 9-33, doi: 10.31992/0869-3617-2023-32-10-9-33 (In Russ., abstract in Eng.)
37. Taheri, R., Nazemi, N., Pennington, S.E., Clark, J.A., Dadgostari, F. (2025). Factors Influencing Educators' AI Adoption: A Grounded Meta-Analysis Review. *Computers and Education: Artificial Intelligence*. No. 9, article no. 100464, doi: 10.1016/j.caeai.2025.100464.
38. Rensfeldt, A.B., Rahm, L. (2023). Automating Teacher Work? A History of the Politics of Automation and Artificial Intelligence in Education. *Postdigital Science and Education*. No. 5, pp. 25-43, doi: 10.1007/s42438-022-00344-x.
39. Gang, D., Yufeng, S., Zhao, Y. (2023) The Innovation of Ideological and Political Education Integrating Artificial Intelligence Big Data with the Support of Wireless Network. *Applied Artificial Intelligence*. Vol. 37, no. 1, article no. 2219943, doi: 10.1080/08839514.2023.2219943.
40. Bulathwela, S., Pírez-Ortiz, M., Holloway, C., Cukurova, M., Shawe-Taylor, J. (2024). Artificial Intelligence Alone Will Not Democratise Education: On Educational Inequality, Techno-Solutionism and Inclusive Tools. *Sustainability*. Vol. 16, no. 2, article no. 781, doi: 10.3390/su16020781.

The paper was submitted 29.04.2026

Accepted for publication 29.05.2026



Пятилетний импакт-фактор
РИНЦ-2023, без самоцитирования

ВОПРОСЫ ОБРАЗОВАНИЯ	3,823
ВЫСШЕЕ ОБРАЗОВАНИЕ В РОССИИ	2,999
ОБРАЗОВАНИЕ И НАУКА	2,979
ПСИХОЛОГИЧЕСКАЯ НАУКА И ОБРАЗОВАНИЕ	2,799
УНИВЕРСИТЕТСКОЕ УПРАВЛЕНИЕ: ПРАКТИКА И АНАЛИЗ	2,075
ИНТЕГРАЦИЯ ОБРАЗОВАНИЯ	1,714
СОЦИОЛОГИЧЕСКИЕ ИССЛЕДОВАНИЯ	1,425
ВОПРОСЫ ФИЛОСОФИИ	0,652
ЭПИСТЕМОЛОГИЯ И ФИЛОСОФИЯ НАУКИ	0,583
ВЫСШЕЕ ОБРАЗОВАНИЕ СЕГОДНЯ	0,531
АЛМА МАТЕР (ВЕСТНИК ВЫСШЕЙ ШКОЛЫ)	0,287
ПЕДАГОГИКА	0,027

Искусственный интеллект и новая система высшего образования: ретроспектива ожиданий в 2023 году и горизонты 2029 года

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-83-101

Резаев Андрей Владимирович – д-р филос. наук, профессор, научный сотрудник¹, профессор кафедры философии², ORCID: 0000-0002-3918-835X, Researcher ID: K-3472-2013, rezaev@hotmail.com

¹ Санкт-Петербургский государственный экономический университет, Санкт-Петербург

Адрес: 191023, г. Санкт-Петербург, наб. канала Грибоедова, д. 30-32, литера А

² Ташкентский государственный экономический университет, Ташкент, Республика Узбекистан

Адрес: 100066, Республика Узбекистан, г. Ташкент, пр-кт Ислама Каримова, д. 49

Трегубова Наталья Дамировна – канд. социол. наук, научный сотрудник¹, доцент кафедры сравнительной социологии², ORCID: 0000-0003-3259-5566, Researcher ID: K-3487-2013, n.tregubova@spbu.ru

Санкт-Петербургский государственный экономический университет, Санкт-Петербург

Адрес: 191023, г. Санкт-Петербург, наб. канала Грибоедова, д. 30-32, литера А

Санкт-Петербургский государственный университет, Санкт-Петербург

Адрес: 191124, г. Санкт-Петербург, ул. Смольного, д. 1/3

***Аннотация.** Три года назад мы опубликовали статью о трансформациях в высшем образовании под влиянием технологий искусственного интеллекта (прежде всего – инструментов типа ChatGPT). Цель настоящей статьи – сверить ожидания 2023 г. с реалиями 2026 г., оценить текущую ситуацию и сформулировать новые прогнозы о состоянии высшего образования в эпоху ИИ. Мы начинаем с описания и анализа результатов контент-анализа публикаций по проблематике ИИ в образовании в ведущих русскоязычных журналах за 2023-2025 гг., что позволяет фиксировать тенденции и проблемы в профессиональной литературе. Затем, на основании проведённого анализа, мы переходим к обсуждению реализованных и нереализованных ожиданий 2023 г., после чего характеризуем риски, с которыми сталкиваются участники высшего образования сегодня. Затем в статье выделяются семь положений о российском высшем образовании в эпоху ИИ, которые могут выступать в качестве источников предположений для дальнейшего научного поиска. В завершение статьи мы формулируем вывод о том, как будет меняться высшая школа в долгосрочной перспективе.*

Ключевые слова: российское высшее образование, трансформация высшего образования, искусственный интеллект, искусственная социальность

Для цитирования: Резаев А.В., Трегубова Н.Д. Искусственный интеллект и новая система высшего образования: ретроспектива ожиданий в 2023 году и горизонты 2029 года // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 83–101. – DOI: 10.31992/0869-3617-2026-35-6-83-101.

Artificial Intelligence and the New Higher Education System: A Retrospective of Expectations in 2023 and Horizons for 2029

Original article

DOI: 10.31992/0869-3617-2026-35-6-83-101

Andrey V. Rezaev – Dr.Sci. (Philosophy), Professor, ORCID: 0000-0002-3918-835X, Researcher ID: K-3472-2013, rezaev@hotmail.com, Researcher¹, Professor of Philosophy Chair²

¹ St. Petersburg State University of Economics, Saint-Petersburg, Russian Federation

Address: 30-32 A, Griboedov canal emb., St. Petersburg, 191023, Russian Federation

² Tashkent State University of Economics, Tashkent, Republic of Uzbekistan

Address: 49 Islam Karimov ave., Tashkent, 100066, Republic of Uzbekistan

Natalia D. Tregubova – Cand.Sci. (Sociology), Researcher¹, Associate Professor, Chair of Comparative Sociology², ORCID: 0000-0003-3259-5566, Researcher ID: K-3487-2013, n.tregubova@spbu.ru

¹ St. Petersburg State University of Economics, Saint-Petersburg, Russian Federation

Address: 30-32 A, Griboedov canal emb., St. Petersburg, 191023, Russian Federation

² St. Petersburg State University, Saint-Petersburg, Russian Federation

Address: 1/3-9, Smolnogo str., Saint-Petersburg, 191124, Russian Federation

Abstract. Three years ago, we published an article examining higher education transformation in the era of artificial intelligence technology expansion. The paper's focus was on LLMs such as ChatGPT.

The purpose of the present paper is to compare the expectations for 2023 with the realities observed in 2026. We aim to assess the current state of higher education and to propose new perspectives for its progress. To begin with, we present and examine the findings from a content analysis of articles about AI in education published in leading Russian-language journals between 2023 and 2025. This helps us to identify prevailing trends and gaps in professional literature. Based on this analysis, we discuss both realized and unrealized expectations for 2023 and subsequently observe the risks currently facing participants in higher education. The article then outlines seven propositions for Russian higher education in the AI era, which may serve as the basis for research questions in future studies. We conclude with a suggestion regarding the long-term transformation of higher education.

Keywords: Russian higher education, transformation of higher education, artificial intelligence, artificial sociality, ChatGPT

Cite as: Rezaev, A.V., Tregubova, N.D. (2026). Artificial Intelligence and the New Higher Education System: A Retrospective of Expectations in 2023 and Horizons for 2029. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 83-101, doi: 10.31992/0869-3617-2026-35-6-83-101 (In Russ., abstract in Eng.).

Введение

В шестом номере тридцать второго тома журнала «Высшее образование в России» мы опубликовали небольшую статью, озаглавленную «ChatGPT и искусственный интеллект в университетах: какое будущее нам ожидать?» [1]. Прошло три года. Скорость изменений в технологиях ИИ колоссальная¹. Изменения в социально-политической жизни, экономических процессах не менее быстрые и значимые. Согласно метрикам журнала, на конец мая 2026 г. статью посмотрели более 3800 человек. Интерес у читателей есть. Мы взяли на себя смелость подготовить новую статью. Целевая ориентация работы – «сверить часы» и определить направления работы тех, кто продолжает служить высоким целям высшего образования.

Мы хотели бы быть более точными и местоимение «нам» в заглавии статьи 2023 г. заменить на «студентам» и «профессорско-преподавательскому составу».

В качестве общего теоретического основания для настоящих рассуждений выступает концепция «искусственной социальности», сформулированная нами в прошлых работах (обоснование концепции и обзор дискуссий см. в: [2]). Мы определяем искусственную социальность как систему взаимодействий с участием ИИ и человека, в которой ИИ может выступать как посредником взаимодействия между людьми, так и самостоятельным участником. Отдельные аспек-

ты искусственной социальности осмысляются в рамках различных теоретических традиций: эффективность информационных обменов – в рамках теорий коммуникации, «сбои» в повседневном взаимодействии – в рамках микросоциологических подходов, склонность к очеловечиванию технологий – в психоаналитической традиции, и т. д. Искусственная социальность в своём развитии выступает основанием для формирования взаимозависимости между людьми и алгоритмами в повседневной жизни общества, а в пределе приводит к формированию взаимозависимости между алгоритмами без участия человека.

Именно это, как представляется, происходит сегодня в высшем образовании. Сначала разработчики создают новые инструменты ИИ, с которыми взаимодействуют отдельные студенты и преподаватели, потом эти взаимодействия становятся всё более распространёнными и типизированными, а затем наступает момент, когда большинство участников образовательного процесса не могут помыслить о том, чтобы обходиться в своей деятельности без алгоритмов. Так было с онлайн-поисковиками, так было с онлайн-переводчиками, сейчас подобная ситуация характеризует взаимодействие с генеративными инструментами ИИ. В пределе возникает ситуация, когда студент пишет курсовую работу, предварительно оценивает её качество и корректирует с помощью ИИ, а преподаватель

¹ В ноябре 2022 года появляется *ChatGPT*, чат-бот, основанный на достаточно «развитой» языковой модели, построенной на архитектуре *GPT (Generative Pre-trained Transformer)*. Суть данной архитектуры и модели заключается в том, что она «предварительно натренирована» на очень большом количестве текстов. На основании данной «тренировки» модель генерирует новые связные, но всё-таки вероятностные ответы на запрос пользователя. То есть модель не ищет готовые фразы, а каждый раз «сочиняет» оригинальный текст, опираясь на выявленные в данных смысловые схемы и контекст диалога с пользователем. Сегодня существуют уже тысячи подобных чат-ботов с гораздо большей вычислительной мощностью и функционалом, чем оригинальная версия.

проверяет эту работу с помощью другого (или такого же) ИИ².

В дальнейших рассуждениях мы будем следовать от частного к общему и от ближнего к дальнему. Мы начнём с систематического обзора публикаций в русскоязычных журналах по проблематике использования ИИ в образовании, затем перейдём к выяснению того, какие ожидания 2023 г. реализовались, какие – нет, и что было нами упущено. После этого укажем на риски, с которыми сталкиваются участники высшего образования в эпоху ИИ. Затем мы выделим семь новых положений о российском высшем образовании в эпоху ИИ, которые могут выступать в качестве оснований концептуальных тезисов и предположений для дальнейших исследований. В завершение статьи мы сформулируем вывод о том, как будет меняться высшая школа в более долгосрочной перспективе.

ИИ в высшем образовании:

на что обращают внимание исследователи

Рассматривая динамику изменений за прошедшие три года, начнём с реакции на эти изменения научного сообщества.

Чтобы отследить отклик исследователей и зафиксировать полученные результаты, мы провели контент-анализ публикаций в ведущих российских журналах. Поиск публикаций производился в феврале 2026 г. в наукометрической базе *eLibrary*. Были выбраны пять наиболее цитируемых журналов в области исследований образования³: 1) «Высшее образование в России» (ВОвР); 2) «Вопросы образования» (ВО); 3) «Психологическая наука и образование» (ПНиО); 4) «Образование и наука» (ОиН); 5) «Универ-

ситетское управление: практика и анализ» (УУПиН).

Для каждого из журналов был произведён поиск публикаций за 2023–2025 гг., имеющих отношение к тематике искусственного интеллекта. Поиск осуществлялся в базе *eLibrary* и, при необходимости просмотра последних номеров за 2025 г., на сайтах журналов. Поисковый запрос выполнялся на русском и английском языках по наличию в названии, аннотации или ключевых словах следующих терминов: искусственный/*artificial*, интеллект/*intellect*, интеллектуальный/*intelligent*, ИИ/*AI*, робот/*robot*, нейросеть / *neural network*, алгоритм/*algorithm*, *ChatGPT*. Результаты поисковой выдачи (описание статей на *eLibrary* и, при необходимости, сами тексты) затем вручную просматривались на соответствие тематике. В выборку включались статьи, в которых использование инструментов ИИ в образовании является основной или одной из основных тем. Всего было обнаружено 62 релевантных статьи.

На графике на *рисунке 1* можно наблюдать стабильный рост интереса авторов к проблематике ИИ в образовании. С 2023 г. по 2025 г. и общее количество статей (шкала слева), и их доля в общем числе публикаций (шкала справа) выросли более чем в три раза.

При этом, как показывает *рисунк 2*, массив публикаций об ИИ распределён неравномерно. Флагманом в отношении развития данной тематики является журнал «Высшее образование в России», где ещё в 2023 г. был опубликован призыв к исследовательскому сообществу обратить внимание на проблему использования ИИ [3]. В данном журнале каждая седьмая публикация попала в нашу

² Студенты под научным руководством одного из авторов настоящей статьи выразили удивление, что их научный руководитель не использует ИИ-инструменты для «быстрого» чтения и проверки их работ. «Вы бы сэкономили много времени» – сказали они. На встречный вопрос: «Вы же сами можете загрузить работы в ИИ, зачем тогда преподаватель?» – студенты, после некоторой заминки, ответили: «С Вами общаться приятнее». Это и есть искусственная социальность на стадии формирования взаимозависимости «человек – алгоритм».

³ Пять журналов из ядра РИНЦ по тематике «Народное образование. Педагогика», лидирующих по числу цитирований из ядра РИНЦ.

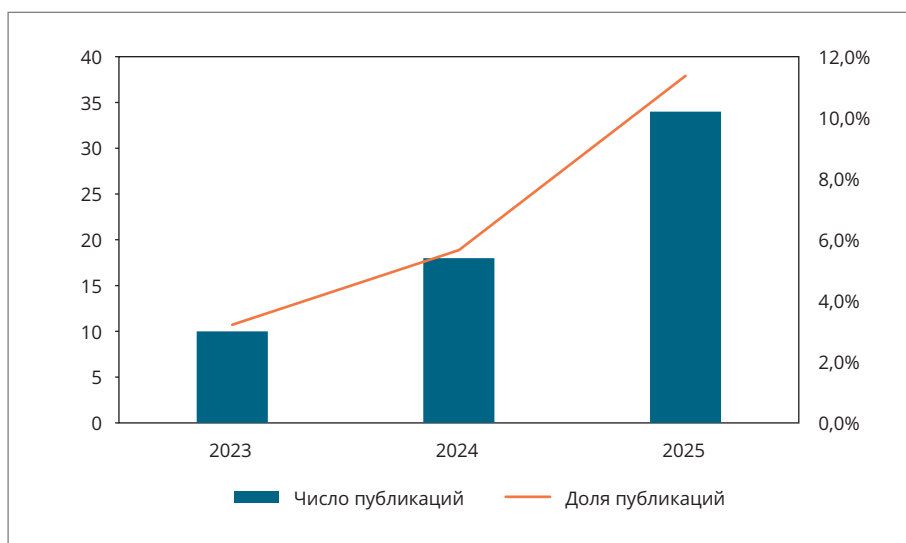


Рис. 1. Динамика публикаций об ИИ в образовании за 2023–2025 гг.

Fig. 1. Dynamics of publications on AI in education for 2023–2025

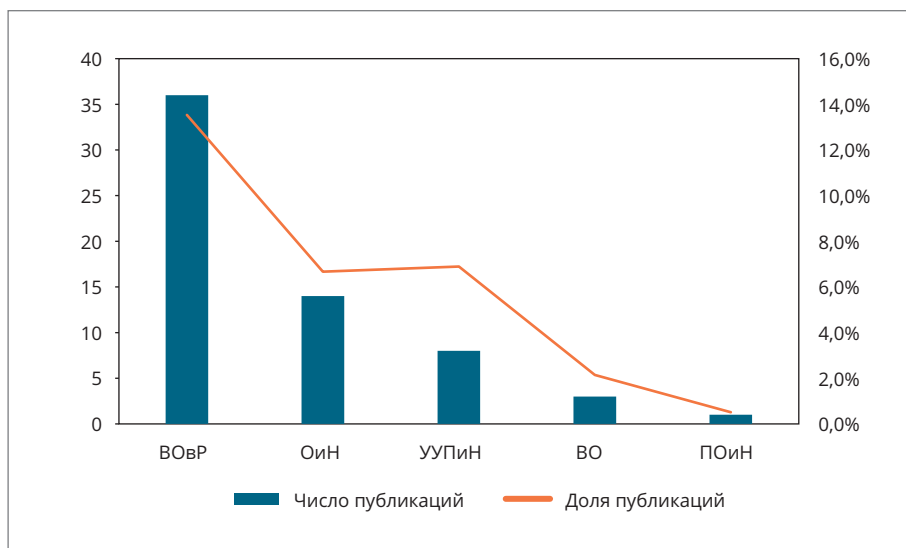


Рис. 2. Распределение статей об ИИ по журналам

Fig. 2. Distribution of AI articles by journal

выборку. Меньшая, но заметная доля публикаций об ИИ характерна для журналов «Образование и наука» и «Университетское управление: практика и анализ». В то же время в «Вопросах образования» и «Психологической науке и образовании» публикации об ИИ появляются эпизодически.

Таким образом, в 2023–2025 гг. мы наблюдаем экспоненциальный рост публикаций об ИИ (каждый год их число практически удваивается) при неравномерно распределённом интересе к данной теме. Примечательно, что интерес в большей степени характерен для журналов, специализирующихся на пробле-

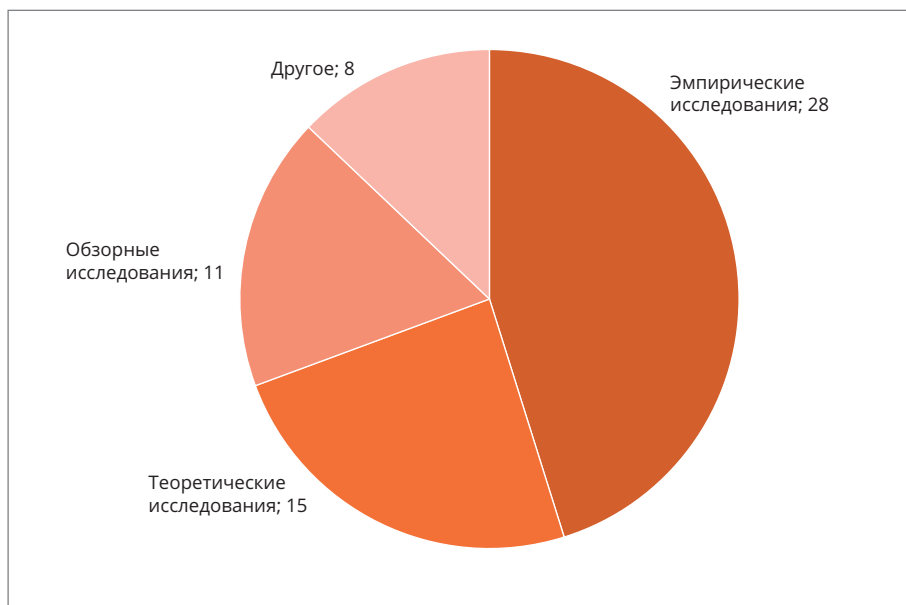


Рис. 3. Распределение статей по исследовательским жанрам
Fig. 3. Distribution of articles by research genres

мах высшего образования. Возможно, это вызвано тем, что исследователи, которые очень часто являются и преподавателями, видят проблемы вхождения ИИ в собственную жизнь и стремятся их осмыслить. Тем более, что работа в университете обеспечивает сравнительно лёгкий доступ к исследовательскому «полю» в сравнении со школами.

Переходя к характеристике содержания публикаций, следует сгруппировать статьи в несколько жанров (Рис. 3):

1) Эмпирические исследования использования ИИ составляют почти половину выборки. Здесь можно выделить такие формы организации исследования как интервью/опрос заинтересованных сторон (студентов, преподавателей, работодателей, выпускников), экспертный опрос, анализ статистических данных, анализ кейсов.

2) Теоретические исследования – второй по популярности жанр: к нему относятся почти четверть статей. Здесь авторы стремятся сформулировать концептуальные основания анализа вхождения технологий ИИ

в образование. Если для эмпирических статей характерны конкретные исследовательские вопросы и ограниченные обобщения, то для теоретических статей более типично рассматривать ИИ как вызов системе образования в целом.

3) Обзор исследований и наработок, как и теоретические исследования, характеризуется стремлением дать общую картину происходящих изменений. Однако здесь в центре внимания оказываются не оригинальные концепции, а систематизация и анализ существующих результатов.

4) Наконец, в выборке представлены иные жанры: описание конкретных разработок, тестирование ИИ как исследовательского инструмента, интервью, приглашение к дискуссии.

Результаты представленных в нашей выборке исследований характеризуют различные аспекты вхождения ИИ в образовательные практики – как в отношении позиций заинтересованных сторон, так и в отношении представленных теоретических перспектив. Вместе с тем знакомство с публикациями

показывает, что социальная аналитика ИИ в российском высшем образовании находится на начальном этапе пути. Дан хороший старт – но уже на этом старте фиксируются несколько проблем.

Первая проблема – разрыв между абстрактными теоретическими рассуждениями о радикальных изменениях и конкретными эмпирическими исследованиями текущих практик. При всей ценности и первых, и вторых, нельзя не видеть пропасть между характеристиками «университета поколения 4.0» [4] или «университета экосистемного типа» [5] – и гораздо более скромными описаниями того, что думают студенты и преподаватели вузов *NN* об использовании моделей ИИ, доступных на первую/вторую половину 202X года (например: [6–8]). Данная проблема ставит перед исследователями вопрос, какие формы организации исследования позволяют хотя бы частично проверить сформулированные концептуальные положения.

Вторая проблема, связанная с первой, состоит в том, что большая часть эмпирических исследований сосредотачивается на фиксации существующих тенденций, без попытки сформулировать научную проблему. Сегодня характер использования инструментов ИИ меняется очень быстро, поэтому не вполне понятно, чью именно «температуру» измеряют чисто описательные исследования, и как долго эта «температура» продержится у конкретного «больного». В качестве контрпримеров – публикации, формулирующих интересные гипотезы и обобщения – можно выделить статью исследовательской группы Я.И. Кузьмина, в которой обсуждается проблема нарастания неравенства при использовании ИИ [9], и исследование Д.П. Ананьина с коллегами, в котором разница в использовании ИИ студентами и преподавателями увязывается с ролью этих инструментов в деятельности первых и вторых [10]. Также следует указать на обзоры, систематизирующие проблемы исследования ИИ в образовании [11; 12].

Третья проблема связана с многоголосицей в теоретических попытках осмысления вхождения ИИ в высшее образование. С одной стороны, значительное число статей (15 теоретических, к которым можно добавить несколько обзорных) стремятся сформулировать, что, как и почему будет менять ИИ в образовании. С другой стороны, между авторами статей практически не возникает диалога. Каждый остаётся со своими тезисами, не критикуя, но и не поддерживая суждения других. Представляется, что это замедляет прогресс в данной области.

Наконец, ещё одна проблема связана с концептуальным смешением, которое имеет место в ряде публикаций. Здесь мы имеем в виду: а) помещение в рубрику «искусственный интеллект» самых разных технологий, без уточнения того, что именно имеется в виду; б) недостаток внимания к тому, как инструменты ИИ влияют на разные уровни образования (прежде всего, школьное и университетское); в) недостаток внимания к разнице между отношением к инструментам ИИ и их реальным использованием.

Зафиксировав ситуацию в исследовательской литературе, перейдём к характеристике изменений в высшем образовании, которые мы можем увидеть сегодня. При этом мы будем стремиться к заполнению пробела между теорией и эмпирическими исследованиями, рассматривая, какие ожидания, зафиксированные в нашей статье 2023 г., оправдались, а какие – нет.

От ожиданий 2023 года к реальности 2026 года

Три года назад, в самом начале первой волны общественного энтузиазма вокруг *ChatGPT* и больших языковых моделей, мы сформулировали несколько теоретико-методологических оснований, на которых можно строить рассуждения о возможных направлениях развития высшей школы в эпоху ИИ [1]. Сегодня, в 2026 г., мы имеем возможность оценить степень реализации

собственных прогнозов и более точно определить векторы дальнейшего развития.

Для верификации прогнозов мы используем двухкомпонентную аналитическую рамку: 1) Сравнительный анализ ожиданий и реализованных изменений; 2) Анализ феноменов, проявившихся трендов и проблем, которые мы не выделили в 2023 г.

Реализованные ожидания

- Тезис о неизбежности внедрения ИИ в образовательный процесс подтвердился полностью. Сегодня генеративные модели используются студентами повсеместно⁴, при этом именно в вузах, наряду с исследовательскими организациями, происходят активные исследования и разработки в области ИИ [14].

- Прогноз о применении ИИ в аккредитации и лицензировании вузов реализовался в большей степени, чем ожидалось. Здесь открываются широкие перспективы для внедрения инструментов ИИ под контролем людей-экспертов [15].

Частично реализованные ожидания

- Предположение об использовании ИИ для адаптации образовательного контента подтвердилось частично. Преподаватели используют ИИ для подготовки материалов курсов, однако эта практика не является повсеместной или всеобщей [6; 10].

- Ожидание роста междисциплинарных исследований и преподавания подтвердилось слабо. Несмотря на то, что модели ИИ действительно облегчают работу с разнородными массивами знаний, институциональные барьеры (кафедральная структура, деканаты, отсутствие межфакультетских механизмов финансирования) оказались сильнее технологических возможностей. В качестве исключения можно выделить несколько флагманских университетов, создавших междисциплинарные лаборатории с активным использованием

ИИ для анализа больших данных (Высшая школа экономики, ИТМО, МГУ).

Нереализованные ожидания

- Интеграция технологий виртуальной и дополненной реальности с ИИ осталась на уровне пилотных проектов. Высокая стоимость оборудования и сложность создания качественного контента ограничили массовое внедрение. Даже в медицинских и технических вузах VR-симуляции с ИИ-поддержкой используются фрагментарно.

- Персонализация учебного процесса с помощью инструментов ИИ также пока не нашла широкого применения в рамках университетов. Далеко не все студенты представляют сущность и потенциал персонализированного обучения, способны выступать субъектами учебного процесса, при этом персонализированное обучение предъявляет новые требования к преподавателям, к которым они не вполне готовы [16].

- Наш прогноз о необходимости определения «табуированных зон», где ИИ не должен использоваться, вызвал дискуссии, но практически не был реализован институционально. Ни Минобрнауки, ни вузовские комитеты по этике не сформулировали чётких регламентов. Результат – не более, чем ситуативные запреты без системной этической рамки.

Что мы не выделили / не предвидели

Во-первых, и это нам представляется принципиальным моментом, мы совсем не предполагали быструю «усталость» от ИИ (*AI fatigue*). К 2025 г. фиксируется феномен снижения интереса к генеративным моделям и инструментам ИИ среди части преподавателей и студентов⁵. После первичного энтузиазма наступило разочарование, связанное с ограниченностью возможностей текущих моделей («галлюцинации», поверхностность анализа) и отсутствием педагогически обо-

⁴ Так, по данным опроса НИУ ВШЭ, 89% студентов российских университетов применяли ИИ для подготовки письменных работ или ответов к экзаменам: см. [13].

⁵ What should the higher education sector do about AI fatigue? – 28.07.2025. – URL: <https://wonkhe.com/blogs/what-should-the-higher-education-sector-do-about-ai-fatigue/> (дата обращения: 10.04.2026).

снованных методик интеграции ИИ в учебный процесс.

Во-вторых, генеративные модели создали проблему, выходящую за рамки классического плагиата. Появились «гибридные тексты», где студент комбинирует сгенерированный моделью контент с собственными фрагментами, что делает традиционные системы антиплагиата неэффективными. Университеты столкнулись с необходимостью переосмысления самого понятия «авторство» [17; 18].

В-третьих, то, что можно назвать «технологическая поляризация». Возник разрыв между вузами-лидерами, создавшими собственные ИИ-лаборатории и интегрировавшими технологии в образовательные программы, и периферийными университетами, где ИИ воспринимается как «модное слово» без реального внедрения. Это несомненно усиливает неравенство в качестве образования, ведёт к различным напряжениям в среде профессорско-преподавательского состава. Кроме того, даже в ведущих вузах возникает разрыв в использовании ИИ студентами, между «технарями» и «гуманитариями», а также между отличниками и отстающими [8].

Таким образом, мы видим, что скорость и направление изменений в высшем образовании за последние три года порождают новые проблемы, многие из которых трудно было предвидеть. Инструменты ИИ (в особенности – генеративного ИИ) становятся частью процесса получения высшего образования, однако это происходит во многом бессистемно, даже на уровне отдельной организации – не говоря о стране в целом, и порождает напряжение – как в отношении необходимости формирования новых компетенций [19; 20], так и в общении между преподавателями и студентами [21; 22]⁶.

⁶ Здесь можно провести параллель с ситуацией в здравоохранении, где основное влияние пока оказывают не специализированные, а универсальные инструменты ИИ-поиска, которые меняют отношения между врачом и пациентом [23].

⁷ См. результаты исследования китайских студентов: [24].

Риски и проблемы:

критический анализ текущего состояния

Масштабное внедрение ИИ-технологий в российское высшее образование сопряжено с комплексом рисков, многие из которых недооценивались три года назад. Для систематизации анализа мы используем классификацию по уровням: индивидуальный (студенты, преподаватели), институциональный (университеты) и системный (вся сфера высшего образования).

Индивидуальные риски

- Риск нераспознавания ошибки. Если вы когда-нибудь что-нибудь создавали, то знаете, что наиболее важная вещь – это даже не идея, а ошибка, которую вы не сразу заметите. Вы будете продолжать что-то делать, но ошибка будет уже внедрена в процесс – и это обойдётся дороже, чем что-либо ещё. Так и с образованием – главный риск в том, что люди не смогут идентифицировать ошибки, которые предлагают модели.

- Риск когнитивной деградации. Чрезмерная опора на ИИ-ассистентов может привести к атрофии критического мышления и навыков самостоятельного анализа⁷.

- Риск деqualификации преподавателей. Если модели ИИ берут на себя рутинные функции (проверка работ, генерация учебных материалов), возникает опасность утраты преподавателями профессиональных компетенций и превращения их в «операторов ИИ-систем».

Институциональные риски

- Здесь следует отметить двойственность в деятельности университетов по внедрению инструментов ИИ. С одной стороны, «риск ничегонеделания», ожидания «команд» сверху. С другой стороны, «риск подмены образования технико-технологическим фетишизмом». Университеты внедряют различные цифровые способы взаимодействия между студентами и преподавателями,

ИИ-модели и инструменты для демонстрации инновационности, не имея ясного понимания педагогической целесообразности. Другими словами, это риск формального использования технологий без улучшения образовательных результатов⁸.

- Следует отметить также и риск зависимости от зарубежных ИИ-инструментов и технологий. Несмотря на развитие российских ИИ-платформ (*YandexGPT*, *GigaChat*), многие передовые инструменты остаются продуктами иностранных компаний (США и Китай). Это создаёт проблемы с суверенитетом данных и устойчивостью образовательных процессов.

Системные риски

- Риск стандартизации мышления. Если большинство студентов используют одни и те же ИИ-модели, обученные на схожих данных, возникает опасность стандартизации образовательных результатов и снижения разнообразия интеллектуальных подходов.

- Риск этической неопределённости. Отсутствие чётких этических регламентов использования ИИ создаёт ситуацию правовой неопределённости. Кто несёт ответственность, если ИИ-ассистент предоставил студенту неверную информацию, использованную в исследовании? Как определить границы допустимого использования ИИ в академических работах?

- Риск дезинформации. ИИ-инструменты могут использоваться для создания убедительных, но фактически неверных материалов. В образовательном контексте это особенно опасно, так как подрывает эпистемологические основы обучения⁹.

Специфические риски в контексте российской реформы высшего образования (2026)

Переход на новую систему высшего образования создаёт дополнительные вызовы для интеграции ИИ-технологий:

- Риск несогласованности технологических и организационных изменений. Одновременное внедрение новых образовательных стандартов и ИИ-инструментов может привести к хаотичности процессов и снижению качества образования в переходный период.

- Риск углубления разрыва между декларируемыми и реальными целями. Если реформа формально провозглашает человеко-ориентированный подход, но фактически стимулирует массовизацию и оптимизацию через ИИ-автоматизацию, возникнет противоречие, деформирующее образовательный процесс.

Высшее образование в эпоху ИИ: ожидания 2026 года

Настоящая статья посвящена оценке изменений, произошедших в российском высшем образовании за три года с момента появления *ChatGPT* и подобных ему генеративных моделей, что позволит зафиксировать некоторые тенденции будущего развития. В продолжение наших рассуждений сформулируем семь положений о дальнейших изменениях в высшем образовании, которые мы надеемся соотнести с реальностью ещё через три года – в 2029 г.

Первое замечание. Высшая школа – наиболее консервативная и наименее подверженная изменениям система, которая,

⁸ Здесь необходимо отметить возможные решения, которые в той или иной степени будут определяться созданием при вузах междисциплинарных комитетов по оценке педагогической эффективности ИИ моделей и инструментов (модель *Evidence-Based AI in Education*). Нам представляется необходимым развивать практики, уже наработанные в Высшей школе экономики, где каждое предложение о внедрении ИИ-технологии проходит экспертизу на соответствие педагогическим целям образовательных программы.

⁹ Здесь также уже существуют некоторые решения – включение в образовательные программы обучения навыкам верификации информации в эпоху ИИ, развитие систем «цифровых следов» (*watermarking*) для ИИ-генерированного контента, позволяющих идентифицировать источник информации.

строго говоря, в своих базовых принципах не изменилась со средних веков (1088 – появление Болонского университета) до становления и развития Болонского процесса (1999 – 2010)¹⁰. Изменения, которые принесла модель *ChatGPT* в высшую школу, стали многократно более принципиальными, по сравнению с тем, что было в период становления Болонского процесса.

Ценности университета формировались в эпоху, когда информация была: а) доступна не всем; б) знания и информация добывались медленно; и в) всё происходило за закрытыми дверями, в «башне из слоновой кости». Причём за закрытыми дверями формировались нерушимые правила игры (формальные и неформальные институты), выстраивались неизменяемые годами структуры (кафедры, факультеты), иерархические бюрократические организационные и управленческие образования, стандартизированные экзаменационные вопросы и тесты, инструкции по получению дипломов и степеней, деятельность вокруг публикационной активности.

Соответственно, базируясь на своих ценностях и бюрократических структурах, университет поставлял обществу: а) доступ к большим и качественным данным, информации и знаниям; б) доступ к «сетям», талантливым людям (студентам и профессорам); в) доступ к дипломам и сертификатам, которые высоко котируются на рынке труда. Эти три фундаментальные ценности были взаимосвязаны и не могли быть получены нигде в другом месте. ИИ принципиально меняет данную ситуацию.

Второе замечание. Университет в эпоху ИИ «соревнуется» не с университетами «нетрадиционного типа» (онлайн-школами или буткемпами). Университеты соревнуются с «фабриками, производящими *intelligence*», даже не с генеративными моделями, а именно с фабриками, которые будут производить в том числе и генеративные модели. И так же,

как было в индустриальную эпоху, фабрики изменяют общество больше, чем школы. Можно ожидать, что технологии ИИ вместе с фабриками по производству информации изменят общество кардинально.

Поэтому в системе образования нельзя продолжать идти по пути *prompt engineering*, то есть подходить к ИИ как к простому инструменту. Точнее – ИИ-грамотность состоит в том, чтобы понять, что ИИ это не просто нечто, решающее конкретную задачу, а нечто, позволяющее в результате решения задачи приносить результат. Реальным дипломом в конце обучения в высшей школе должна стать «демонстрация результатов». Подобно тому, как курсант военного училища должен продемонстрировать, как он умеет летать или стрелять. Подобно тому, как спортсмен демонстрирует результаты тренировок в забеге, прыжке. По-видимому, очень правильно будет присмотреться, как организовано обучение в театральных вузах. Каким образом организуется мотивация будущих актёров и тренирующихся спортсменов на получение результатов.

Главным становится не выписка из зачётной книжки с оценками за экзамены и зачёты, а количество реализованных проектов и участие в решении задач из практики конкретной жизни.

Третье замечание. Процессы изменения, которые уже сегодня начались в университетах, будут напоминать те процессы, с которыми столкнулись газеты, а затем телевизионные компании. Когда газеты не заметили приход цифровых медиа. Когда телекомпании не увидели реального конкурента в платформах *VK* и *YouTube*. Сначала а) игнорирование наступления новой реальности, затем б) шоковое состояние от потери бывшего влияния на умы и действия студентов, потом в) паника, следом г) попытки выстроить новые схемы под старые бренды, после д) консолидация лучших сил для решения задач информирования в условиях ИИ, и наконец – ж)

¹⁰ В общем виде о том, что принёс Болонский процесс, см. в [25].

предложение новых систем и структурных образований.

Что могут сейчас делать университеты для того, чтобы выжить? Они могут/будут делать то, что делал бизнес, когда появился Интернет и когда началась цифровизация экономики. Они начнут и будут продолжать вводить курсы по «ИИ-грамотности». Они разрешат студентам использовать генеративные модели в процессе обучения и преподавателям для подготовки к занятиям. Они будут развивать новые методы «борьбы» с обманом и плагиатом. Другими словами, в ближайшее время, университеты будут рассматривать ИИ в системе высшего образования как проблему, которую можно решить традиционными для университета способами. То есть, университеты будут приспосабливать ИИ «под себя», под свои структуры и институты¹¹.

Но маловероятно, что университеты будут успешными, развиваясь в этом направлении. Скорее всего, надо будет подстраивать структуры и институты университетов «под ИИ», под новые фабрики производства информации. Реальность существования/выживания университетов требует осознания того, что человечество, с одной стороны, переходит к новым типам взаимоотношения между людьми [2]. С другой стороны, процессы обучения, рынок труда в новой реальности будут развиваться в ситуации, где информация доступна всем очень быстро и может быть персонализирована. Реальность для структур, институтов и процессов образования становится такой, что образование перестаёт быть строго регламентированным

с точки зрения времени и места его получения, оно становится адаптивным (*adaptive education*).

Четвёртое замечание. Задача университета на завтра состоит не в том, чтобы бравировать прошлыми заслугами. Студентам нужна будет атмосфера сотрудничества, постоянного поиска нового, возможности отстаивать собственную точку зрения на основе критического самостоятельного анализа информации, предоставляемой ИИ-моделями. Цель университета в том, чтобы реально вырастить личность, а не дать студенту бумагу о том, что он личность, потому что прослушал те или иные курсы и отучился пять лет в том или ином университете.

Изменится то, каким образом рассуждают исследователи «экономики образования». Сегодня существует то, что называется *Sheepskin Effect* («эффект диплома»): доход резко возрастает именно в момент получения формального диплома или степени, а не просто с каждым дополнительным годом обучения¹². Человек с почти завершённым обучением (например, не хватает одного «кредита» до бакалавра) зарабатывает заметно меньше, чем тот, кто формально получил степень, хотя их объём знаний почти одинаков. Здесь также следует привести аргумент экономиста Брайана Каплана, который утверждает, что значительная часть ценности образования на рынке труда связана не с реальными знаниями и навыками, а с тем, что диплом подаёт работодателю сигнал о качествах выпускника (способности, дисциплина, конформизм) [26].

¹¹ Так, когда *ChatGPT* получил широкое распространение в студенческой среде, университеты во всём мире стали а) (или) запрещать использование ИИ для использования студентами, а потом/или б) применять другие ИИ-модели, для того, чтобы обнаружить, что студенты используют ИИ для домашних работ, и наказывать студентов за это. То есть, в качестве образца поведения с новыми ИИ-технологиями использовались старые практики в обучении – не пускать или контролировать. Другими словами, продолжать тренировать студентов на запоминание чего-то, что может быть когда-то использовано. Нам представляется, что сейчас и в будущем следует поступать прямо наоборот. Если выпускник будет применять ИИ в своей работе, то почему ему не научиться работать с ИИ-моделями в период прохождения курсов в университете?

¹² Название связано с тем, что дипломы традиционно печатали на овечьей коже (*sheep skin*).

ИИ может предоставить другие способы, определяющие сигналы, которые демонстрировали не быстрые и не дешёвые университетские дипломы. Причём ИИ может определить те качества человека, которые нужны потенциальному работодателю очень быстро, дешево и в самых конкретных условиях. Можно ожидать, что востребованность навыков, связанных с работой с ИИ – как и развитие самих инструментов ИИ для оценки качества полученного образования – приведёт к уменьшению «эффекта диплома», к необходимости, помимо диплома, иметь иные подтверждения ИИ-компетентности.

И здесь возникает необходимость определиться: какие навыки нужны студентам и выпускникам в эпоху ИИ? Среди таких навыков несомненно должны быть: а) гибкость / пластичность / адаптированность к изменяющимся реалиям, б) креативность, в) критичность по отношению к получаемой информации, г) критичность в отношении принимаемых решений.

Что здесь можно получить от ИИ? 1) ИИ как тьютор (инструктор), который не отменяет преподавателя, но помогает ему структурировать, как и что делать; 2) ИИ как редактор текстов, которые пишет студент; 3) ИИ как *ТА* (*teaching assistant*); 4) ИИ как помощник в конкретных практических заданиях. Сегодня проблема состоит в том, чтобы понять, что есть *AI-literacy* («ИИ-грамотность»).

Пятое замечание. Университетам следует обратить внимание на то, что ИИ не сможет делать.

ИИ не сможет отличать истинную информацию от ложной.

ИИ не сможет (а скорее, сможет навредить) развивать сотрудничество самых различных специалистов при решении задач. Это должно стать целью университетского образования. То, что называлось междисциплинарностью.

ИИ не сможет помочь выработать то, что называется «вкусом», то есть понимание

того, что есть красота, «что такое хорошо», в чём значимость того или иного действия/продукта для людей, этические составляющие человеческого труда.

ИИ не сможет решить такие проблемы выпускников как: «Я долго учился и наконец получил диплом, но не уверен, как применить свои знания на практике», «Мне сложно связать полученные знания со своими реальными проектами на полученной работе или будущими карьерными целями».

ИИ-модели в принципе не могут генерировать цели человеческой деятельности и повседневной жизни, они смогут очень быстро предоставить варианты, статистические выкладки, но не цель. ИИ сможет предоставить в лучшем случае усреднённый ответ (в худшем случае – вариант манипуляции от производителя той или иной модели) о том, зачем человек живёт, трудится, для кого и для чего нужна всякая человеческая активность.

Шестое замечание. ИИ-модели сегодня приносят замещение взаимодействий между студентами и преподавателями, студентами и студентами. Уходит общение, приходит обмен информацией. Происходит «взаимодействие с помощью больших пальцев» на смартфоне, а не при помощи голоса, обмена взглядами и невербальных форм общения.

Суть дела в том, что ИИ изменяет коренным образом саму ситуацию получения информации. Преподаватели вузов не выдержат конкуренции по четырём направлениям: а) ИИ чётко видит, где каждый студент испытывает затруднение при получении и переработки информации и сможет каждому объяснить снова и снова любые затруднения; б) причём ИИ всегда может буквально моментально (без необходимости прочитать новую книгу или статью) давать новые объяснения, опять же каждому студенту; в) ИИ может быть доступен студенту всегда, в любое время дня и ночи; и, что, возможно, наиболее существенно, г) ИИ всегда будет повторять студенту необходимую информацию, до тех пор, пока студенту будут требоваться по-

вторные разъяснения, не испытывая никаких эмоциональных проблем с повторением.

Нельзя не учитывать, насколько эмоциональная/человеческая составляющая образования влияет на студента (которому стыдно признаться, что он не понял, в чём-то не разобрался) и на преподавателя (который не в состоянии повторять одно и то же без конца в аудитории, где много студентов). Строго говоря, революция, которую несёт ИИ в процесс обучения, касается не только (и не столько) содержания обучения, сколько уверенности студента, что не понимать – не стыдно, и уверенности, что у каждого студента может быть свой персонализированный инструктор, который «видит», что один студент может хорошо понимать математику, но не философию, а другой, наоборот, видит философскую составляющую проблемы, но с трудом воспринимает необходимость делать вычисления.

ИИ изменяет центральный компонент сегодняшнего образовательного процесса в высшей школе. Не класс, не курсовая группа, а индивидуальный студент становится центром. Каждый студент использует ту или иную генеративную модель для получения информации. Лекция и занятие в аудитории становятся самым медленным активом, который не только забирает время и энергию, но и не может предоставить персонализированные варианты получения информации и знания каждому студенту в отдельности, без студенческой группы.

При этом ИИ оставляет за скобками тот факт, что в университетах во время обучения происходит социализация студентов, общение с себе подобными и с теми, кто знает и умеет больше, чем ты. В университете «происходит» человеческая жизнь, а не просто активное/пассивное существование в условиях искусственной социальности.

И здесь возникают перспективы для серьёзных трансформаций роли преподавате-

ля, который должен служить не транслятором знаний, а фигурой, которая привлекает и удерживает внимание студентов. Понятно, что эти роли со-существовали в высшем образовании на протяжении веков. Однако сегодня именно вторая, дополнительная роль выходит на первый план, в то время как передача информации всё в большей степени монополизирована инструментами ИИ. Преподаватель не может продолжать быть инструктором, хорошо исполняющим инструкцию, он должен стать «профессором» в подлинном значении этого слова в русском языке.

Профессор в эпоху ИИ не есть проверяющий, что студент *не* знает, *не* смог запомнить, *не* смог выучить, *не* сделал¹³. Профессор является личностью, который генерирует вместе со студентом идеи на протяжении семестра, показывает, как концентрировать внимание, формирует культуру постановки вопросов и осмысление необходимости их решения. Он должен оценить, какую задачу студент не сможет решить, показать каковы траектории его дальнейшего развития. Профессор должен разбираться в очень сложном процессе: как измерять на экзаменах то, как студент «мыслит», а не красиво пишет или делает красивые презентации с помощью ИИ; как измерить способности рациональных рассуждений, а не форматировать то, что предоставляет ИИ-модель. Он должен оценивать варианты ответов на возникающие проблемы в условиях неопределённости, а не конкретные ответы на predetermined вопросы. Главное, что должно оцениваться, – как быстро студент может понять, что он «не прав», предлагая решение той или иной проблемы, каковы параметры его устойчивости к «неверным ходам», как быстро студент может предложить иное решение, как слаженно он работает в группе.

Седьмое (и последнее) замечание. Мы продолжаем утверждать, что высшее обра-

¹³ «Большинство учителей тратят время, спрашивая, чего не знает ученик, вместо того чтобы выяснить, на что он способен» (А. Эйнштейн).

зование должно стать новой областью знания, причём такой областью, которая будет ориентирована на междисциплинарность (равно как и ИИ) и методологический плюрализм. Причём добавятся сюда не только языковые модели, но и «большие мировые модели» (*large world models*) или роботы. Разница между роботами и языковыми моделями в том, что роботы не могут тренироваться на словах. Им нужен опыт. Высшее образование – это тренировка не только на словах, даже с помощью языковых моделей, но и на опыте, эксперименте, в процессе реального творения/делания.

Заключение

Анализ, проведённый в данной работе, показал, что исследования вхождения ИИ в практики высшего образования всё ещё находятся на начальном этапе, при этом динамика текущих изменений такова, что очень трудно прогнозировать что-то даже в пределах двух-трёх лет. В рамках статьи мы попытались соотнести текущие тенденции и ожидания, чтобы зафиксировать наиболее важные векторы трансформации.

В заключение попробуем сформулировать более долгосрочный «прогноз» о том, как высшее образование будет меняться с развитием технологий ИИ. Здесь мы не берём на себя риск определять конкретные сроки – только общее направление развития.

Высшее образование дня сегодняшнего не помогает *решить* проблемы завтрашнего дня, скорее именно высшее образование в том виде, в котором оно сейчас существует, *создаёт* эти проблемы. Сегодня университетский мир (феодално-иерархичный), который ориентирован на фиксированные и медленные/времяёмкие экзаменационные процедуры, на комиссии, приходит в противоречие с быстро меняющимися и легко совпадающими с рыночными требованиями реалиями участия ИИ в повседневной жизни, в обучении и производстве. Модель, предлагаемая университетом – «сначала обучение (время- и ресурсозатратное) – потом рабо-

та», сталкивается с моделью, которую предлагает ИИ: «работа и обучение проходят одновременно», ИИ выступает постоянным помощником и в работе, и в натаскивании на решение тех задач, которые требуют работа.

Высшее образование перестаёт быть «этапом/ступенью жизни», но становится непрерывным процессом повышения квалификации в условиях искусственной социальности.

По сути дела, ИИ не просто помогает студентам учиться. ИИ изменяет матрицу и суть получения информации, на основе которой происходит формирование мировоззрения. Обучение перестаёт быть простым следованием инструкциям и воспроизведением ответов, которые преподаватель/инструктор ожидает от студентов и может их оценить. Студенты-отличники, которые становились отличниками потому, что научились существовать в высшей школе в системе фиксированных учебных планов, конкретных экзаменационных вопросов, чётких инструкций по получению аттестатов и дипломов, будут с трудом адаптироваться к новым реалиям, которые приходят вместе с ИИ.

Если алгоритм может легко сформулировать ответы на различные вопросы, то задача образования – научить задавать вопросы, а не воспроизводить ответы на прошлые вопросы. Задача высшего образования состоит в том, чтобы выпускник университета строил свои решения, основываясь на умении анализировать сложные и постоянно меняющиеся ситуации в условиях нестабильности, на основе самостоятельно сформулированных вопросов, решение которых приводит к реальным прорывам в той или иной сфере.

Литература

1. Резаев А.В., Трегубова Н.Д. ChatGPT и искусственный интеллект в университетах: какое будущее нам ожидать? // Высшее образование в России. – 2023. – Т. 32, № 6. – С. 19–37. – DOI: 10.3192/0869-3617-2023-32-6-19-37.
2. Резаев А.В., Трегубова Н.Д. Иная социальность: опыт теоретической рефлексии.

- сии // Социологические исследования. – 2025. – № 11. – С. 13–24. – DOI: 10.7868/S3034577X25110021.
3. Ивахненко Е.Н., Никольский В.С. ChatGPT в высшем образовании и науке: угроза или ценный ресурс? // Высшее образование в России. – 2023. – Т. 32, № 4. – С. 9–22. – DOI: 10.31992/0869-3617-2023-32-4-9-22.
 4. Ефимов В.С., Лаптева А.В. Поколения университетов: особенности культивируемых типов мышления. Каким будет мышление в Университете 4.0? // Высшее образование в России. – 2024. – Т. 33, № 8-9. – С. 95–122. – DOI: 10.31992/0869-3617-2024-33-8-9-95-122.
 5. Матковская Я. С., Русаева Е. Ю. Диалектика становления университетов экосистемного типа и её влияние на формирование специалистов будущего // Университетское управление: практика и анализ. – 2023. – Т. 27, № 3. – С. 95–114. – DOI: 10.15826/umpra.2023.03.026.
 6. Буякова К.И., Дмитриев Я.А., Иванова А.С., Фещенко А.В., Яковлева К.И. Отношение студентов и преподавателей к использованию инструментов с искусственным интеллектом в вузе // Образование и наука. – 2024. – Т. 26, № 7. – С. 160–193. – DOI: 10.17853/1994-5639-2024-7-160-193.
 7. Сысоев П.В. Искусственный интеллект в образовании: осведомлённость, готовность и практика применения преподавателями высшей школы технологий искусственного интеллекта в профессиональной деятельности // Высшее образование в России. – 2023. – Т. 32, № 10. – С. 9–33. – DOI: 10.31992/0869-3617-2023-32-10-9-33.
 8. Тихонова Н.В., Поморцева Н.П. Выпускная квалификационная работа в вузе в условиях распространения искусственного интеллекта: взгляд студентов // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 112–135. – DOI: 10.31992/0869-3617-2025-34-6-112-135.
 9. Кузьминов Я.И., Кручинская Е.В., Груздев И.А., Наумов А.А. Отстающие и опережающие: как студенты используют генеративный искусственный интеллект в образовательных целях // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 9–35. – DOI: 10.31992/0869-3617-2025-34-6-9-35.
 10. Ананин Д.П., Комаров Р.В., Реморенко И.М. «Когда честно – хорошо, для имитации – плохо»: стратегии использования генеративного искусственного интеллекта в российском вузе // Высшее образование в России. – 2025. – Т. 34, № 2. – С. 31–50. – DOI: 10.31992/0869-3617-2025-34-2-31-50.
 11. Катаев Д.В., Беляев Д.А., Тарасов А.Н. Динамика научного дискурса об искусственном интеллекте в образовании: библиометрический анализ и тематическое моделирование // Высшее образование в России. – 2025. – Т. 34, № 11. – С. 145–168. – DOI: 10.31992/0869-3617-2025-34-11-145-168.
 12. Кошкина Е.А., Бордовская Н.В., Гнедых Д.С., Хромова М.А., Демьянчук Р.В., Исакова М.П., Балышев П.А. Генеративный искусственный интеллект в высшем образовании: обзор теоретических подходов и практик применения // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 36–57. – DOI: 10.31992/0869-3617-2025-34-6-36-57.
 13. Кырма В.В., Орлова Е.О., Розмаинский И.В. Эмпирический анализ чатинга среди студентов разных университетов // Journal of Institutional Studies. – 2025. – Т. 17, № 3. – С. 136–156. – DOI: 10.17835/2076-6297.2025.17.3.136-156.
 14. Искусственный интеллект в России: разработка и применение / П.Б. Рудник (рук. авт. кол.), В.А. Абашкин, Г.И. Абдрахманова и др.; под ред. А.М. Гохберга, П.Б. Рудника, Г.И. Абдрахмановой. – Москва: ИСИЭЗ ВШЭ, 2025. – URL: <https://issek.hse.ru/mirror/pubs/share/1053986532.pdf> (дата обращения: 10.04.2026).
 15. Болотов В.А., Мотова Г.Н. Ассистент или судья: использование искусственного интеллекта в аккредитации образовательных программ // Университетское управление: практика и анализ. – 2025. – Т. 29, № 4. – С. 5–16. – DOI: 10.15826/umpra.2025.04.027.
 16. Сысоев П.В. Персонализированное обучение на основе технологий искусственного интеллекта: насколько готовы современные студенты к новым возможностям получения образования // Высшее образование в России. – 2025. – Т. 34, № 2. – С. 51–71. – DOI: 10.31992/0869-3617-2025-34-2-51-71.
 17. Сысоев П.В. Этика и ИИ-плагиат в академической среде: понимание студентами вопросов соблюдения авторской этики и проблемы плагиата в процессе взаимодействия с генеративным искусственным интеллектом // Высшее образование в России. – 2024. – Т. 33, № 2. – С. 31–53. – DOI: 10.31992/0869-3617-2024-33-2-31-53.

18. Земцов Д.И., Груздев И.А. «Цифровой кентавр»: совместное обучение человека и ИИ в университете // Высшее образование в России. – 2025. – Т. 34, № 10. – С. 47–62. – DOI: 10.31992/0869-3617-2025-34-10-47-62.
19. Пузанова Ж.В., Ларина Т.И. Интеграция прикладного искусственного интеллекта в магистерские программы непрофильных направлений: вызовы, тренды и перспективы // Высшее образование в России. – 2025. – Т. 34, № 8-9. – С. 33–53. – DOI: 10.31992/0869-3617-2025-34-8-9-33-53.
20. Конколь М.М., Марьина Е.Д. Методологические основания системы метацифровой компетентности (на примере языкового образования) // Образование и наука. – 2025. – Т. 27, № 9. – С. 9–29. – DOI: 10.17853/1994-5639-2025-9-9-29.
21. Никольский В.С. Коммуникативный искусственный интеллект: концептуализация новой реальности в образовании // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 152–168. – DOI: 10.31992/0869-3617-2025-34-6-152-168.
22. Резаев А.В., Трегубова Н.Д. Внедрение инструментов искусственного интеллекта в сферу высшего образования: взгляд с позиций социально-институциональной парадигмы общения // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 80–90. – DOI: 10.31992/0869-3617-2025-34-6-80-90.
23. Резаев А.В., Трегубова Н.Д., Иванова А.А. Трансформация медицинской экспертизы в России в условиях искусственной социальности: опыт сравнительного анализа // Социологическое обозрение. – 2026. – Т. 25, № 1. – С. 177–296. – DOI: 10.17323/1728-192x-2026-1-177-206.
24. Tian J., Zhang R. Learners' AI dependence and critical thinking: The psychological mechanism of fatigue and the social buffering role of AI literacy // Acta Psychologica. – 2025. – Vol. 260. – Article no. 105725. – DOI: 10.1016/j.actpsy.2025.105725.
25. Rezaev A.V. Bologna Process: On the way to a Common European Higher Education Area // Editors: Penelope Peterson, Eva Baker, Barry McGaw, International Encyclopedia of Education (Third Edition). – Elsevier, 2010. – P. 772–778. – DOI: 10.1016/B978-0-08-044894-7.00169-X.
26. Caplan B. The Case Against Education: Why the Education System Is a Waste of Time and Money. – Princeton: Princeton University Press, 2018. – ISBN: 978-0691174655.

Благодарности. Исследование выполнено за счёт гранта Российского научного фонда № 26-18-00243.

Статья поступила в редакцию 15.04.2026

Принята к публикации 29.05.2026

References

1. Rezaev, A.V., Tregubova, N.D. (2023). ChatGPT and AI in the Universities: An Introduction to the Near Future. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 6, pp. 19-37, doi: 10.31992/0869-3617-2023-32-6-19-37 (In Russ., abstract in Eng.).
2. Rezaev, A.V., Tregubova, N.D. (2025). Alternate Sociality: Theoretical Reflections. *Sotsiologicheskie issledovaniya = Sociological Studies*. No. 11., pp. 13-24, doi: 10.7868/S3034577X25110021 (In Russ., abstract in Eng.).
3. Ivakhnenko, E.N., Nikolskiy, V.S. (2023). ChatGPT in Higher Education and Science: a Threat or a Valuable Resource? *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 4, pp. 9-22, doi: 10.31992/0869-3617-2023-32-4-9-22 (In Russ., abstract in Eng.).
4. Efimov, V.S., Lapteva, A.V. (2024). Generations of Universities: Features of Cultivated Thinking Types. What Will the Thinking be Like in University 4.0? *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 33, no. 8-9, pp. 95-122, doi: 10.31992/0869-3617-2024-33-8-9-95-122 (In Russ., abstract in Eng.).
5. Matkovskaya, Ya.S., Rusyaeva, E.Yu. (2023). Ecosystem-Type Universities' Formation Dialectics and Its Influence Over Future Specialists. *University Management: Practice and Analysis*. Vol. 27, no. 3, pp. 95-114, doi: 10.15826/umpa.2023.03.026 (In Russ., abstract in Eng.).

6. Buyakova, K.I., Dmitriev, Ya.A., Ivanova, A.S., Feshchenko, A.V., Yakovleva, K.I. (2024). Students' and Teachers' Attitudes Towards the Use of Tools with Generative Artificial Intelligence at the University. *The Education and science journal*. Vol. 26, no. 7, pp. 160-193, doi: 10.17853/1994-5639-2024-7-160-193 (In Russ., abstract in Eng.).
7. Sysoyev, P.V. (2023). Artificial Intelligence in Education: Awareness, Readiness and Practice of Using Artificial Intelligence Technologies in Professional Activities by University Faculty. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 32, no. 10, pp. 9-33, doi: 10.31992/0869-3617-2023-32-10-9-33 (In Russ., abstract in Eng.).
8. Tikhonova, N.V., Pomortseva, N.P. (2025). Final Qualification Paper in University in the Context of Artificial Intelligence Proliferation: University Students' Perspective. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 112-135, doi: 10.31992/0869-3617-2025-34-6-112-135 (In Russ., abstract in Eng.).
9. Kuzminov, Ya.I., Kruchinskaia, E.V., Gruzdev, I.A., Naumov, A.A. (2025). Falling Behind and Getting Ahead: Student Use of Generative AI in Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 9-35, doi: 10.31992/0869-3617-2025-34-6-9-35 (In Russ., abstract in Eng.).
10. Ananin, D.P., Komarov R.V., Remorenko, I.M. (2025). "When Honesty is Good, for Imitation is Bad": Strategies for Using Generative Artificial Intelligence in Russian Higher Education Institutions. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 2, pp. 31-50, doi: 10.31992/0869-3617-2025-34-2-31-50 (In Russ., abstract in Eng.).
11. Kataev, D.V., Belyaev, D.A., Tarasov, A.N. (2025). Dynamics of Scientific Discourse on Artificial Intelligence in Education: Bibliometric Analysis and Thematic Modeling. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 11, pp. 145-168, doi: 10.31992/0869-3617-2025-34-11-145-168 (In Russ., abstract in Eng.).
12. Koshkina, E.A., Bordovskaya, N.V., Gnedykh, D.S., Khromova, M.A., Demyanchuk, R.V., Iskharova, M.P., Balyshv, P.A. (2025). Generative Artificial Intelligence in Higher Education: A Review of Theoretical Approaches and Application Practices. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 36-57, doi: 10.31992/0869-3617-2025-34-6-36-57 (In Russ., abstract in Eng.).
13. Kyrma, V.V., Orlova, E.O., Rozmainsky, I.V. (2025). Empirical Analysis of Cheating among Students from Different Universities. *Journal of Institutional Studies*. Vol. 17, no. 3, pp. 136-156, doi: 10.17835/2076-6297.2025.17.3.136-156 (In Russ., abstract in Eng.).
14. Rudnik, P.B., et al. (2025). *Artificial Intelligence in Russia: Developments and Applications*. Moscow: ISSEK HSE. Available at: <https://issek.hse.ru/mirror/pubs/share/1053986532.pdf> (accessed 10.04.2026). (In Russ.).
15. Bolotov, V. A., Motova, G. N. (2025). Assistant or Judge: The Role of Artificial Intelligence in Study Program Accreditation. *University Management: Practice and Analysis*. Vol. 29, no. 4, pp. 5-16, doi: 10.15826/umpa.2025.04.027 (In Russ., abstract in Eng.).
16. Sysoyev, P.V. (2025). Personalized Learning Based on Artificial Intelligence: How Ready Are Modern Students for New Educational Opportunities. *Vysshee obrazovanie v Rossii= Higher Education in Russia*. Vol. 34, no. 2, pp. 51-71, doi: 10.31992/0869-3617-2025-34-2-51-71 (In Russ., abstract in Eng.).
17. Sysoyev, P.V. (2024). Ethics and AI-Plagiarism in an Academic Environment: Students' Understanding of Compliance with Author's Ethics and the Problem of Plagiarism in the Process of Interaction with Generative Artificial Intelligence. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 33, no. 2, pp. 31-53, doi: 10.31992/0869-3617-2024-33-2-31-53 (In Russ., abstract in Eng.).
18. Zemtsov, D.I., Gruzdev, I.A. (2025). "Digital Centaur": Human-AI Collaborative Learning

- in Universities. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 10, pp. 47-62, doi: 10.31992/0869-3617-2025-34-10-47-62 (In Russ., abstract in Eng.).
19. Puzanova, Zh.V., Larina, T.I. (2025). Integrating Applied Artificial Intelligence into Non-Technical Master's Programs: Conceptual Models and Practical Approaches. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 8-9, pp. 33-53, doi: 10.31992/0869-3617-2025-34-8-9-33-53 (In Russ., abstract in Eng.).
20. Konkol, M.M., Marina, E.D. (2026). Methodological Foundations of the Meta-Digital Competence System: A Case Study in Language Education. *Obrazovanie i nauka = The Education and Science Journal*. Vol. 27, no. 9, pp. 9-29, doi: 10.17853/1994-5639-2025-9-9-29 (In Russ., abstract in Eng.).
21. Nikolskiy, V.S. (2025). Communicative Artificial Intelligence: Conceptualizing a New Reality in Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 152-168, doi: 10.31992/0869-3617-2025-34-6-152-168 (In Russ., abstract in Eng.).
22. Rezaev, A.V., Tregubova, N.D. (2025). AI Tools in Higher Education through the Lens of the Social-Institutional Paradigm of Social Intercourse. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 80-90, doi: 10.31992/0869-3617-2025-34-6-80-90 (In Russ., abstract in Eng.).
23. Rezaev, A.V., Tregubova, N.D., Ivanova, A.A. (2026). A Comparative Analysis of the Medical Expertise Transformation in Russia in the Age of Artificial Sociality. *Russian Sociological Review*. Vol. 25, no. 1, pp. 177-296, doi: 10.17323/1728-192x-2026-1-177-206 (In Russ., abstract in Eng.).
24. Tian, J., Zhang, R. (2025). Learners' AI Dependence and Critical Thinking: The Psychological Mechanism of Fatigue and the Social Buffering Role of AI literacy. *Acta Psychologica*. Vol. 260, article no. 105725, doi: 10.1016/j.actpsy.2025.105725.
25. Rezaev, A.V. (2010). Bologna Process: On the Way to a Common European Higher Education Area. In: P. Peterson, E. Baker, B. McGaw (Eds.) *International Encyclopedia of Education* (Third Edition). Elsevier. Pp. 772-778, doi: 10.1016/B978-0-08-044894-7.00169-X.
26. Caplan, B. (2018). *The Case Against Education: Why the Education System Is a Waste of Time and Money*. Princeton: Princeton University Press. ISBN: 978-0691174655.

Acknowledgement. The research was supported by the Russian Science Foundation, project No. 26-18-00243.

The paper was submitted 15.04.2026

Accepted for publication 29.05.2026

Тревожность относительно искусственного интеллекта в образовательной среде: психометрическая адаптация шкалы на выборке российских студентов

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-102-126

Максименко Александр Александрович – д-р социол. наук, канд. психол. наук, доцент, профессор факультета социальных наук, SPIN-код: 7449-3003, ORCID: 0000-0003-0891-4950, Scopus ID: 57219144362, Researcher ID: AAZ-8789-2021, maximenko.al@gmail.com

Национальный исследовательский университет «Высшая школа экономики», Москва, Россия
Адрес: 109000, г. Москва, ул. Мясницкая, д. 20

Круз-Карденас Хорхе – PhD в области экономики и управления бизнесом, руководитель исследовательского центра ESTec факультета экономики, управления и бизнеса, ORCID: 0000-0002-4575-6229, Scopus ID: 56000985700, Researcher ID: PUF-4615-2026, jorgecruz@uti.edu.ec

Технологический университет Индоамерики, Кито, Эквадор

Адрес: 170103, г. Кито, ул. Мачала, 5 км, д. 1/2

Забелина Екатерина Вячеславовна – д-р психол. наук, доцент, профессор кафедры психологии Института образования и практической психологии, ORCID: 0000-0002-2071-6466, Scopus ID: 633966, Researcher ID: V-1905-2018, katya_k@mail.ru

Челябинский государственный университет, Москва, Россия

Адрес: 454001, г. Челябинск, ул. Братьев Кашириных, д. 129

Курапов Сергей Владимирович – канд. социол. наук, старший научный сотрудник, SPIN-код: 5004-1747, ORCID: 0000-0002-2003-7439, v_ksv@mail.ru

Российский научно-исследовательский институт культурного и природного наследия имени Д.С. Лихачева, Москва, Россия

Адрес: 129366, г. Москва, ул. Космонавтов, д. 2

Певнева Анжела Николаевна – канд. психол. наук, заведующая кафедрой общей и социальной психологии факультета психологии, ORCID: 0000-0001-6304-0649, Scopus ID: 57212474012, Researcher ID: PGF-8373-2026, pevneva_AN@grsu.by

Гродненский государственный университет имени Янки Купалы, Гродно, Беларусь

Адрес: 230005, г. Гродно, ул. Гаспадарчая, д. 23

Аннотация. Сегодня в образовательном пространстве можно наблюдать явление тревожности по поводу искусственного интеллекта, которое в значительной степени влияет на готовность принимать технологические решения и планировать своё профессиональное будущее. В эмпирических исследованиях ИИ-тревожность определяется как сложный феномен, включающий спектр негативных эмоциональных реакций, возникающих в связи с развитием и потенциальным влиянием искусственного интеллекта на различные сферы жизни. Учитывая влияние ИИ-тревожности на обучение, карьерные выборы, участие в цифровой экономике и психическое благополучие, возникает потребность в её целостном концептуальном описании и надёжной диагностике. Цель исследования – представить надёжный и валидный русскоязычный инструмент для диагностики ИИ-тревожности у студентов. Принятая за основу и модифицированная для студенческой выборки (N = 521) «Шкала ИИ-тревожности» (М. Алиайбани с соавторами) показала приемлемый уровень внешней и внутренней валидности, а также надёжности по внутренней согласованности. Конфирматорный факторный анализ подтвердил пяти-факторную структуру опросника: тревога, связанная с точностью и надёжностью ИИ (ассигасу); опасения, связанные с академической честностью (plagiarism); тревога, обусловленная нормативной неопределённостью (guidelines), опасения утраты учебных и когнитивных навыков (skill loss), а также снижение учебной мотивации (motivation). Внешняя валидность подтверждается наличием корреляций с нейротизмом и добросовестностью как чертами личности, образом будущего, отношением к искусственному интеллекту, а также с некоторыми социально-демографическими характеристиками. Адаптированный и модифицированный психодиагностический инструмент для изучения ИИ-тревожности открывает возможности для исследований на русскоязычной студенческой выборке.

Ключевые слова: искусственный интеллект, ИИ-тревожность, российские студенты, шкала ИИ-тревожности, академическая добросовестность, неопределённость, самооэффективность, учебная мотивация, психометрическая адаптация

Для цитирования: Максименко А.А., Круз-Карденес Х., Забелина Е.В., Курапов С.В., Певнева А.Н. Тревожность относительно искусственного интеллекта в образовательной среде: психометрическая адаптация шкалы на выборке российских студентов // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 102–126. – DOI: 10.31992/0869-3617-2026-35-6-102-126

Artificial Intelligence Anxiety in the Educational Context: Psychometric Adaptation of the Scale on a Sample of Russian Students

Original article

DOI: 10.31992/0869-3617-2026-35-6-102-126

Alexander A. Maksimenko – Dr.Sci. (Sociology), Professor, Department of Social Sciences, SPIN-code: 7449-3003, ORCID: 0000-0003-0891-4950, maximenko.al@gmail.com

National Research University Higher School of Economics University, Moscow, Russian Federation
Address: 20 Myasnetskaya str., Moscow, 109000, Russian Federation

Jorge Cruz-Cardenas – PhD in Economics and Business Management, Head of the ESTec Research Center, Faculty of Economics, Management and Business, ORCID: 0000-0002-4575-6229, jorgecruz@uti.edu.ec
Indoamerica Technological University, Quito, Ecuador
Address: bldg. 1/2, Machala str., 5 km, Quito, 170103, Ecuador

Ekaterina V. Zabelina – Dr.Sci. (Psychology), Associate Professor, Professor of the Department of Psychology at the Institute of Education and Practical Psychology, ORCID: 0000-0002-2071-6466, katya_k@mail.ru
Chelyabinsk State University, Chelyabinsk, Russian Federation
Address: 129 Brat'yev Kashirinykh str., Chelyabinsk, 454001, Russian Federation

Sergey V. Kurapov – Candidate of Sciences in Sociology, Senior Researcher, SPIN-code: 5004-1747, ORCID: 0000-0002-2003-7439, v_ksv@mail.ru
Russian Research Institute for Cultural and Natural Heritage named after D.S. Likhachev, Moscow, Russian Federation
Address: 2 Kosmonavtov str., Moscow, 129366, Russian Federation

Angela N. Pevneva – Candidate of Sciences in Psychology, Head of the Department of General and Social Psychology at the Faculty of Psychology, ORCID: 0000-0001-6304-0649, Researcher ID: PGF-8373-2026, pevneva.angela@rambler.ru
Yanka Kupala State University of Grodno, Grodno, Belarus
Address: 23 Gasparadchaya str., Grodno, 230005, Belarus

Abstract. Today, there is a growing anxiety about artificial intelligence in the educational space, which greatly affects the willingness to make technological decisions and plan for one's professional future. In empirical studies, AI anxiety is defined as a complex phenomenon including a range of negative emotional responses that shapes in people due to development and potential impact of artificial intelligence on various spheres of life. Given the impact of AI anxiety on learning, career choices, participation in the digital economy, and mental well-being, there is a need for its holistic conceptual description and reliable diagnosis. The aim of the study is to present a reliable and valid Russian-language tool for diagnosing AI anxiety in students. The adopted basis and modified for a student sample (N = 521) "Artificial Intelligence Anxiety Scale" (M.H. Alkhaibani et al.) showed an acceptable level of the external and internal validity as well as reliability on internal consistency. The factor analysis confirmed the five-factor structure of the questionnaire: concerns related to AI accuracy and reliability; concerns related to academic integrity (plagiarism); and concerns stemming from regulatory uncertainty (guidelines), fears of loss of learning and cognitive skills, as well as a decrease in motivation. External validity is confirmed by the correlations with neuroticism and conscientiousness as personality traits, image of the future, attitude towards artificial intelligence, and some socio-demographic characteristics. Adapted and modified psychodiagnostic tool for studying AI anxiety opens up opportunities for research on a Russian-speaking student sample.

Keywords: artificial intelligence, AI anxiety, Russian students, AI anxiety scale, academic integrity, uncertainty, self-efficacy, educational motivation, psychometric adaptation

Cite as: Maksimenko, A.A., Cruz-Cardenes, J., Zabelina, E.V., Kurapov, S.V., Pevneva, A.N. (2026). Artificial Intelligence Anxiety in the Educational Context: Psychometric Adaptation of the Scale on a Sample of Russian Students. *Vyshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 102-126, doi: 10.31992/0869-3617-2026-35-6-102-126 (In Russ., abstract in Eng.).

Введение

Искусственный интеллект (ИИ) перестал быть исключительно прикладным алгоритмом и стал важным символом ближайшего будущего: экономического, профессионального и даже экзистенциального [1]. В научном дискурсе обсуждение ИИ сопровождается не только надеждами на рост эффективности и трансформацию отраслей [2] и качества жизни, но и довольно широким спектром опасений: от потери рабочих мест до гипотетического «конца человеческой эры» [3]. Эти страхи, фрагментарно присутствующие в средствах массовой информации, закрепляются в индивидуальном опыте и формируют специфическую форму тревоги, связанную непосредственно с ИИ, но не с технологиями в целом. Систематические обзоры подчёркивают, что тревожность относительно ИИ (ИИ-тревожность) уже стала значимым фактором, влияющим на готовность людей принимать технологические решения и планировать своё профессиональное будущее [4; 5].

ИИ-тревожность распределяется неравномерно между группами населения и способна усиливать существующие неравенства. Менее образованные работники, представители уязвимых профессий и жители стран с более хрупкими системами социальной защиты чаще рассматривают ИИ как угрозу своему человеческому капиталу, а не как ресурс развития [4; 6]. В профессиональных контекстах ИИ-тревожность влияет на готовность к организационным изменениям, инновационное поведение и даже безопасность труда [7; 8]. Таким образом, речь идёт не просто о субъективном страхе, а о факторе, опосредующем траектории включения людей в экономику будущего.

Наряду с этим ИИ-тревожность имеет ярко выраженное отношение к будущему в экзистенциальном плане. В ряде работ подчёркивается, что угрозы, приписываемые ИИ, затрагивают базовые буферы тревоги (*anxiety buffers*) – самоуважение, чувство предсказуемости мира и значимость меж-

личностных привязанностей [1; 9]. Участники исследований описывают страх не только перед автоматизацией труда, но и перед возможной девальвацией человеческого опыта и компетентности, что особенно отчётливо звучит в нарративах специалистов, столкнувшихся с риском вытеснения из-за ИИ [10; 11]. Тем самым ИИ-тревожность оказывается связана с тем, как люди осмысливают своё место в долгосрочном будущем, переживая его либо как пространство утраты, либо как ресурс для адаптации.

С учётом того, что ИИ-тревожность влияет на обучение, карьерные выборы, участие в цифровой экономике и психическое благополучие, возникает потребность в её целостном концептуальном описании. Современные работы (от разработки шкал до феноменологических исследований) пытаются не только зафиксировать уровень тревоги, но и показать, какие психологические механизмы, социально-демографические характеристики и представления о будущем структурируют это явление [5; 12; 13].

Определение и концептуализация ИИ-тревожности

В эмпирических исследованиях ИИ-тревожность определяется как многомерный конструкт, описывающий спектр негативных эмоциональных реакций: от беспокойства и нервозности до выраженного страха, которые возникают у человека в связи с развитием, внедрением и потенциальным влиянием систем искусственного интеллекта на разные сферы жизни. Одна из наиболее влиятельных концептуализаций предложена при разработке Шкалы ИИ-тревожности (*Artificial Intelligence Anxiety Scale*), где выделяются четыре измерения: тревога обучения, тревога замещения рабочей силы, социотехническая слепота и тревога конфигурации ИИ [12]. Более поздние работы, применяющие и адаптирующие эту шкалу, подтверждают устойчивость четырехфакторной структуры ИИ-тревожности у студентов, учителей и представителей широкой общественности,

одновременно демонстрируя вариативность относительной значимости отдельных измерений в разных группах [14–16].

Вместе с тем ряд исследователей расширяет понимание ИИ-тревожности в результате включения дополнительных типов страхов, связанных с приватностью, предвзятостью алгоритмов, непрозрачностью и этическими рисками. Так, при анализе страха перед ИИ через призму технологических возможностей выделяются страх нарушения приватности, тревога по поводу алгоритмической предвзятости, экзистенциальная тревога и опасения, связанные с непрозрачностью сложных систем [13]. В сфере автономного транспорта страхи касаются не только технической надёжности, но и утраты контроля и неуверенности в том, как ИИ будет оценивать человеческие жизни в конфликтных ситуациях [17].

Экзистенциальный компонент выделяется в исследованиях страха перед ИИ как функции приписываемых ему когнитивных и эмоциональных способностей: по мере того как участники начинают воспринимать ИИ как квазисубъект, обладающий сознанием или эмоциями, возрастает ощущение радикальной угрозы человеческому статусу [18]. Через перспективу теории управления страхом смерти показывается, что сценарии «экстремальной автоматизации» подрывают мировоззренческие и ценностные опоры, усиливая тревожность и способствуя защите групповой идентичности [1; 9].

Учитывая разнообразие этих подходов – от узкого понимания как специфической технологической тревожности до широкой экзистенциальной конструкции – *под рабочей дефиницией ИИ-тревожности* понимается устойчивое, когнитивно и эмоционально структурированное состояние ожидания негативных последствий, связанных с использованием или распространением систем искусственного интеллекта, направленное в будущее и включающее опасения относительно занятости, социального статуса, смысла и предсказуемости собственной жизни и общества. Такое состояние фор-

мируется в зависимости от индивидуально-психологических особенностей человека, его социально-демографических характеристик и уровня цифровых навыков.

Психологические факторы ИИ-тревожности

Одним из наиболее систематически описанных факторов ИИ-тревожности выступает технологическая готовность. Исследование на основе модели технологической готовности (*Technology Readiness*) среди студентов показало, что тормозящие компоненты технологической готовности (склонность к дискомфорту, скепсис, ощущение уязвимости) положительно связаны со всеми измерениями ИИ-тревожности: страхом обучения, тревогой замещения, социотехнической слепотой и тревогой конфигурации [19]. Напротив, «позитивные» компоненты технологической готовности оказываются связанными лишь с меньшей социотехнической слепотой и не защищают от других видов тревоги. Похожие результаты получены в выборках общего населения и сотрудников: низкая частота использования компьютеров и слабое знание об ИИ ассоциируются с более негативными установками и повышенной тревожностью [20; 21].

Доминирующую позицию в формировании ИИ-тревожности играют самоэффективность и психологическая «наделённость полномочиями». В контексте человеческо-компьютерного взаимодействия показано, что самоэффективность в обучении технологиям опосредует влияние работы с ИИ на тревогу, связанную с освоением этих технологий, одновременно модифицируя связь между стрессорами и принятием ИИ [22; 23]. В среде маркетологов высокая самоэффективность в отношении ИИ сглаживает негативное влияние тревоги приватности, предвзятости и непрозрачности на намерение использовать генеративные модели [24]. Психологическая «воодушевленность» и ощущение компетентности связаны со снижением тревоги обучения, но парадоксально могут усиливать на-

сторожённость в отношении потенциального вреда ИИ, то есть сопровождаться более критической рефлексией [25; 26].

Личностные черты формируют устойчивый фон ИИ-тревожности. Исследование с использованием модели «Большой пятёрки» показывает положительное влияние нейротизма на общий уровень ИИ-тревожности и все её субшкалы, а также отрицательную связь с позитивным отношением к ИИ [20]. Сочетание высокого уровня нейротизма и выраженной ИИ-тревожности ассоциируется с отвержением ИИ как угрозы устойчивому экономическому развитию. Нетерпимость к неопределённости выступает специфическим механизмом, опосредующим связь духовного/экзистенциального интеллекта и ИИ-тревожности: при высокой нетерпимости к неопределённости духовные ресурсы не снижают страх, а, напротив, усиливают его, особенно в вопросах обучения и замещения рабочих мест [27].

Социально-демографические характеристики, связанные с ИИ-тревожностью

Гендер оказывается одним из наиболее устойчивых предикторов ИИ-тревожности. В ряде работ женщины демонстрируют более высокие уровни ИИ-тревожности, меньшее доверие к ИИ и более низкую самооценку знаний и навыков использования технологий [26; 28]. Адаптация шкалы ИИ-тревоги на турецкой выборке показывает, что женщины-учителя набирают более высокие баллы по сравнению с мужчинами [29; 30]. Исследования в области автономного транспорта и медицинского ИИ дополнительно фиксируют, что женщины чаще выражают опасения в связи с безопасностью, приватностью и возможной дискриминацией [1; 17; 31]. В то же время при высоких уровнях ИИ-тревоги гендерные различия в её выраженности сглаживаются, что позволяет интерпретировать гендерные эффекты как отражение различий в технологической социализации и доступе к цифровым ресур-

сам, а не принципиально разных психологических реакций [28].

Возрастные различия проявляются менее линейно. С одной стороны, респонденты более старшего возраста чаще сообщают о социотехнической слепоте и ощущении отставания от технологического развития [5]. С другой стороны, опасения по поводу потери работы и экономической нестабильности выражены как у молодых, так и у старших возрастных групп, но мотивируются разными образами будущего. Молодые люди боятся неуспешного входа на рынок труда, более взрослые – утраты накопленного человеческого капитала [6; 32].

Уровень образования и профессиональный статус вносят заметный вклад в дифференциацию ИИ-тревожности. Более низкий уровень образования последовательно ассоциируется с повышенной тревогой по поводу обучения ИИ и замещения рабочих мест, что отражает переживание слабой защищённости на рынке труда [5; 14]. Среди студентов выявлена связь между ИИ-тревогой и тревогой по поводу будущей безработицы: чем слабее академические позиции и чем менее определённой кажется профессиональная траектория, тем выше оба вида тревоги [33]. Профессиональные группы также различаются: будущие педагоги по технологии демонстрируют меньшую ИИ-тревожность, чем будущие преподаватели математики и естественных наук, вероятно, потому что в их профессиональной идентичности технологический аспект изначально встроен [30].

На макроуровне кросс-национальные сравнения показывают значительные различия между странами: более высокий средний уровень ИИ-страхов фиксируется в Индии, Саудовской Аравии и США и более низкий – в Турции, Японии и Китае [6]. Эти различия связаны не только с экономическим развитием, но и с культурными представлениями о профессиях, историческим опытом автоматизации и образами ИИ в публичном пространстве. Подобные вариации подчёркиваются также в работах, анализирующих

общественные тревоги относительно медицинского ИИ и религиозных контекстов использования технологий [1; 34].

Несмотря на актуальность проблемы ИИ-тревожности, её влияние на социальные, психологические и экзистенциальные аспекты жизни людей в мире, её исследования в российском контексте представлены дефицитарно. Особенно перспективным видится исследование ИИ тревожности в образовательной среде среди студентов вузов. Однако для этого необходимы валидные и надёжные методики, диагностирующие тревожность в отношении ИИ. Цель этого исследования – адаптировать шкалу ИИ-тревожности на выборке российских студентов.

Методы исследования

Для диагностики тревожности, связанной с использованием искусственного интеллекта в образовательной среде, был выбран 20-пунктный опросник «Шкала ИИ-тревожности» (*Artificial Intelligence Anxiety Scale*), первоначально разработанный М.Х. Альшайбани, В.М. Аль-Рахми, Р.М. Тавафак для академического контекста [35]. Выбор данной методики обусловлен её многомерной структурой, позволяющей дифференцировать различные аспекты ИИ-тревожности, что соответствует современным представлениям о данном конструкте как о сложном психологическом феномене. Изначально методика разрабатывалась для научной среды, вследствие чего была модифицирована под задачи исследования.

С целью проверки внешней валидности применялись следующие инструменты.

1. Краткий пятифакторный опросник личности (*Ten-Item Personality Inventory, TIPI*) в адаптации [36] содержит 10 утверждений, изучающих такие черты личности, как экстраверсия, доброжелательность, добросовестность, нейротизм и открытость опыту.

2. Отношение к ИИ у студентов. Методика состоит из 14 пунктов, объединённых в две шкалы: Утилитарно-позитивное отношение к ИИ и Озабоченность негативными последствиями использования.

3. Шкала «Тёмное будущее» [37] в адаптации Т.А. Нестика и А.А. Журавлева [38]. Шкала представляет собой скрининговый вариант методики З. Залеского «Тревога по поводу будущего», состоящий из 5 утверждений.

4. Семантический дифференциал «Наше будущее» [39]. В основе методики – приём семантического дифференциала, часто используемый для изучения психологического времени. Согласно рекомендациям по обработке, был проведён факторный анализ утверждений, который показал два чётких фактора, объясняющих 54,93% общей дисперсии. Первый фактор, образованный такими дефинициями будущего, как «наполненное», «моё», «прекрасное», «светлое», «осмысленное», «открытое к изменениям», «близкое», «значительное», «исполненное надежд», был обозначен как «экзистенциальное будущее». Второй фактор, включающий определения «предсказуемое», «безопасное», «отчётливое», «зависит от меня», «зависит от нас», был назван «управляемое будущее».

В исследовании также регистрировались социодемографические переменные: пол, возраст, уровень урбанизации и субъективный уровень дохода.

Расчёты проводились с использованием программного обеспечения STATA 19.

Выборка

Исследование проводилось в начале декабря 2025 года с использованием метода онлайн-анкетирования. Объём выборки составил 521 респондент, 69,5% ($n = 362$) – юноши, и 30,5% ($n = 159$) – девушки¹. Возраст участников варьировался от 17 до 23 лет,

¹ Преобладание юношей в выборке (69,5%) объясняется тем, что исследование проводилось на базе технических вузов, где студенты мужского пола традиционно составляют большинство на IT-специальностях, предполагающих активное взаимодействие с инструментами ИИ.

средний возраст составил 18,6 года. Выборка формировалась методом критериального невероятностного отбора среди студентов российских вузов. В качестве респондентов выступили студенты организаций высшего образования. Участие в опросе было добровольным и анонимным; перед заполнением анкеты все участники давали информированное согласие на обработку данных. Критериями включения в выборку являлись: статус студента вуза, наличие опыта использования инструментов искусственного интеллекта в учебной деятельности (в т. ч. вайбкодинг), реализуемого диалогом с ИИ (не менее одного раза за последний месяц), и согласие на участие в исследовании.

По уровню урбанизации выборка характеризовалась следующим распределением: большинство респондентов проживали в мегаполисах (исключая Москву) – 58,5%; в областных центрах – 19,2%; в районных центрах – 18,4%; наименьшую долю составили жители Москвы – 3,8%. Субъективная оценка материального положения показала, что половина респондентов (50,5%) отнесли свой доход к среднему уровню. Почти треть (29,0%) оценили его как низкий, а 12,9% назвали себя «сводящими концы с концами». Высокий и очень высокий доход отметили 6,0% и 1,7% участников соответственно, что в сумме составило 7,7%.

Результаты исследования

Поскольку исходный инструмент был разработан для оценки тревожности у исследователей и отражал специфику научной деятельности, необходимо было не только провести его языковую адаптацию, но и содержательную модификацию, направленную на приведение формулировок в соответствие с образовательным контекстом и опытом студентов. Сначала применялась процедура прямого и обратного перевода: утверждения были переведены с английского языка на русский с учётом особенностей учебной деятельности и типичных ситуаций использования ИИ студентами. Далее неза-

висимым переводчиком был выполнен обратный перевод на английский язык, после чего проводилось сопоставление исходной и полученной версий с целью выявления и устранения возможных смысловых расхождений. Дополнительно адаптированная версия прошла экспертную оценку с участием специалистов в области психологии и лингвистики, что позволило обеспечить содержательную и концептуальную эквивалентность инструмента.

Содержательная модификация опросника заключалась, во-первых, в переработке формулировок утверждений с ориентацией на студенческую аудиторию, а во-вторых, в расширении структуры шкалы за счёт включения дополнительных утверждений. В частности, к исходным 20 утверждениям были добавлены 4 новых пункта (21–24), отражающие влияние ИИ на учебную мотивацию и поведенческие аспекты деятельности. Таким образом, итоговая версия опросника включала 24 утверждения и предполагала пятифакторную структуру.

Для выявления латентной структуры инструмента был проведён разведочный факторный анализ с использованием метода главных компонент. Полученные результаты позволили выделить пять факторов, со значениями, превышающими 1 [40], соответствующих как исходной теоретической модели, так и дополнительному мотивационному компоненту (Табл. 1).

В этом первоначальном решении фактор 1 был наиболее релевантным. Как видно из таблицы 1, собственная величина 6,62 демонстрирует сильную объяснительную силу фактора 1 [41], объясняющего 28% общей дисперсии. Однако это первоначальное решение было сложно в интерпретации. Для получения более лёгкого интерпретируемого решения был проведён факторный анализ с косоугольным вращением, следуя процедуре, описанной М. Алшайбани с соавторами [35]. Результаты показаны в таблице 2. Кроме того, в таблице 2 отображается значение уникальности, то есть отклонение элемента

Таблица 1

Первоначальные факторы и их значения

Table 1

Relevant factors of the initial solution

Фактор	Собственное значение	Доля дисперсии
Фактор 1	6,62	0,28
Фактор 2	2,81	0,12
Фактор 3	2,23	0,09
Фактор 4	1,46	0,06
Фактор 5	1,27	0,05

Таблица 2

Результаты эксплораторного факторного анализа после вращения

Table 2

Exploratory Factor Analysis Results with Rotation

Переменные	Фактор 1	Фактор 2	Фактор 3	Фактор 4	Фактор 5	Уникальность
Q1_1	-0,002	-0,127	-0,92	0,209	0,849	0,333
Q1_2	-0,091	0,078	0,354	-0,353	0,449	0,552
Q1_3	0,044	-0,095	0,004	0,095	0,794	0,388
Q1_4	-0,041	0,350	-0,138	0,123	0,534	0,487
Q1_5	0,001	0,841	-0,145	0,146	-0,103	0,363
Q1_6	0,095	0,476	0,061	-0,103	0,238	0,563
Q1_7	0,003	0,711	0,024	-0,012	0,153	0,375
Q1_8	0,109	0,405	0,130	-0,137	0,214	0,614
Q1_9	0,039	0,864	-0,123	0,144	-0,192	0,330
Q1_10	0,722	-0,295	0,027	0,131	0,115	0,459
Q1_11	0,789	0,069	-0,010	-0,002	0,018	0,334
Q1_12	0,864	0,123	-0,063	-0,080	-0,039	0,239
Q1_13	0,799	0,037	0,046	-0,039	-0,004	0,325
Q1_14	0,757	-0,080	-0,088	0,147	0,016	0,440
Q1_15	0,604	0,253	0,214	-0,167	-0,069	0,374
Q1_16	-0,063	0,120	0,177	0,637	0,161	0,401
Q1_17	-0,062	-0,034	-0,001	0,844	0,151	0,282
Q1_18	-0,048	0,151	0,341	0,562	-0,007	0,385
Q1_19	0,051	0,028	0,068	0,790	0,075	0,293
Q1_20	0,057	0,251	0,173	0,580	-0,080	0,417
Q1_21	-0,006	-0,056	0,676	0,297	-0,023	0,377
Q1_22	0,011	-0,390	0,788	0,024	-0,010	0,452
Q1_23	0,016	0,096	0,667	0,151	-0,058	0,418
Q1_24	0,002	0,054	0,691	0,188	-0,168	0,405

или показателя, которое не объясняется полученным факторным решением [41].

Таблица 2 показывает, что пункты чётко сгруппированы по 5 факторам на основе факторных нагрузок, равных или превышающих 0,4 [41]. Это решение также позволило

обозначить 5 факторов в соответствии с оригинальной шкалой, предложенной М. Алшайбани с соавторами [35], и измерением, добавленным в этом исследовании.

В таблице 3 представлены характеристики 5 факторов, полученных с помощью

Таблица 3

Названные факторы и их характеристики

Table 3

Labeled factors and characteristics

Фактор	Название	Доля дисперсии	Альфа Кронбаха
Фактор 1	Руководство	0,19	0,87
Фактор 2	Плагиат	0,18	0,78
Фактор 3	Мотивация	0,17	0,75
Фактор 4	Потеря навыков	0,16	0,86
Фактор 5	Точность	0,12	0,66

метода главных компонент с косоугольным вращением, а также процентная доля общей дисперсии, объяснённая каждым фактором. При этом фактор 1 (руководство) объяснил 19% общей дисперсии, что немного выше других факторов. В таблице 3 также представлены коэффициенты альфа Кронбаха как мера надёжности шкал по внутренней согласованности. Дж. Хайр с коллегами [40] предполагают, что значения, превышающие 0,7, приемлемы в подтверждающем исследовании, а значения, превышающие 0,6, приемлемы в поисковом исследовании. Как видно из таблицы 3, значения коэффициента альфа Кронбаха были удовлетворительными для всех факторов, за исключением фактора 5 (точность). Для улучшения значения этого коэффициента, было проведено удаление элемента, однако это не улучшило значительно результат. Это является ограничением и заслуживает последующих проверок на других выборках.

Первый фактор (руководство) представлен утверждениями 10–15, описывающими недостаточную определённую нормативных рамок использования ИИ: отсутствие чётких разъяснений со стороны университетов, единых требований преподавателей, прозрачных правил и официальных рекомендаций. Данный фактор фиксирует институциональный источник тревожности, связанный с неопределённостью регулирующих норм. Второй фактор (плагиат) включает утверждения 5–9, в которых фиксируются опасения относительно возможного

нарушения академической добросовестности, возникновения плагиата, генерации неоригинального контента и изменения стандартов выполнения учебных заданий. Этот фактор отражает тревогу, обусловленную неясностью границ допустимого использования ИИ и возможными санкциями со стороны образовательной системы. Третий фактор (мотивация), введённый в рамках модификации опросника, включает утверждения 21–24, касающиеся снижения мотивации к самостоятельной работе, прокрастинации, снижения уверенности в собственных знаниях и уменьшения вовлеченности в учебный процесс. Выделение данного фактора указывает на то, что ИИ-тревожность затрагивает не только когнитивные и нормативные компоненты, но и мотивационно-регулятивные компоненты учебной деятельности. Четвёртый фактор (утрата навыков) объединяет утверждения 16–20, включая высказывания о том, что чрезмерное использование ИИ может ухудшать учебные навыки, снижать способность к критическому мышлению, препятствовать развитию самостоятельности, приводить к утрате аналитических навыков и уменьшать необходимость в самостоятельной подготовке. Данный фактор отражает представление студентов о потенциальной деградации когнитивных и учебных ресурсов в условиях активного использования ИИ. Пятый фактор (точность) объединяет утверждения 1–4, отражающие сомнения студентов в достоверности, интерпретируемости и академической релевантности ре-

Результаты конфирматорного анализа: стандартизированные коэффициенты

Таблица 4

Table 4

CFA Standardized Results

	Коэффициент	Ошибка	z	$P > z $
ИИ-тревожность (AI_Anxiety)->Точность (Accuracy)	0,54	0,08	7,09	0,000
ИИ-тревожность (AI_Anxiety)->Плагиат (Plagiarism)	0,80	0,06	12,72	0,000
ИИ-тревожность (AI_Anxiety)->Руководство (Guidelines)	0,58	0,05	10,86	0,000
ИИ-тревожность (AI_Anxiety)->Утрата навыков (Skill_loss)	0,55	0,05	10,70	0,000
ИИ-тревожность (AI_Anxiety)->Мотивация (Motivation)	0,51	0,06	8,22	0,000
cov (Accuracy-Plagiarism)	0,31	0,12	2,56	0,010
cov (Skill_loss-Motivation)	0,63	0,05	13,67	0,000

зультатов, получаемых с помощью ИИ. Он описывает когнитивный аспект тревожности, связанный с недоверием к технологическим системам.

Следующий этап анализа заключался в оценке полученной модели с использованием конфирматорного анализа. Названия, предложенные в таблице 3, использовались для обозначения латентных переменных – факторов. Была представлена модель второго порядка, включающая два уровня латентных переменных (всего 6). На первом уровне находились пять латентных переменных: точность, плагиат, руководство и потеря навыков [35], плюс латентная переменная «мотивация», добавленная в этом исследовании. На следующем уровне модели была добавлена переменная ИИ-тревожность.

Модель, подлежащая тестированию (Рис.1), была построена на основе следующих предполагаемых взаимосвязей:

1. Тревога по поводу искусственного интеллекта определяет обеспокоенность точностью, плагиатом, отсутствием руководящих принципов, потерей навыков и снижением мотивации.

2. Каждая из этих пяти скрытых переменных определяет, как люди отвечают на соответствующие утверждения (наблюдаемые переменные) в анкете.

3. Интуитивно ожидается, что сомнения в точности и опасность плагиата будут связаны. Аналогично, ожидается, что потеря

навыков и снижение мотивации будут коррелировать.

Сочетание индексов пригодности для модели ($Chi^2(245) = 767,22, p < 0,001$; $RMSEA = 0,064$; $CFI = 0,892$; $SRMR = 0,060$) является приемлемым [40; 42]. Хи-квадрат должен быть незначительным, но обычно он не используется как индекс пригодности в CFA из-за его высокой чувствительности к характеристикам выборки [40]. В этом отношении $RMSEA$, CFI и $SRMR$ считаются более надёжными [40]. Результаты модели CFA показаны в таблице 4. Все пути в модели имеют значительные стандартизированные коэффициенты (Табл. 4), что подтверждает гипотетическую структуру взаимосвязи. Согласно стандартизированным коэффициентам, наиболее сильная связь в модели была между ИИ-тревожностью и плагиатом. Следовательно, можно сделать вывод, что главной проблемой для студентов при использовании ИИ является плагиат. Полная модель со стандартизированными оценками представлена на рисунке 1.

В рамках проверки внешней валидности был проведён корреляционный анализ (r -Пирсона) со шкалами, изучающими отношение к будущему. Анализ выявил интерпретируемые взаимосвязи с некоторыми субшкалами ИИ-тревожности у студентов. Показатель «Плагиат» отрицательно взаимосвязан со шкалами «экзистенциальное будущее» ($r = -,101$; $p = ,022$) и «управляемое

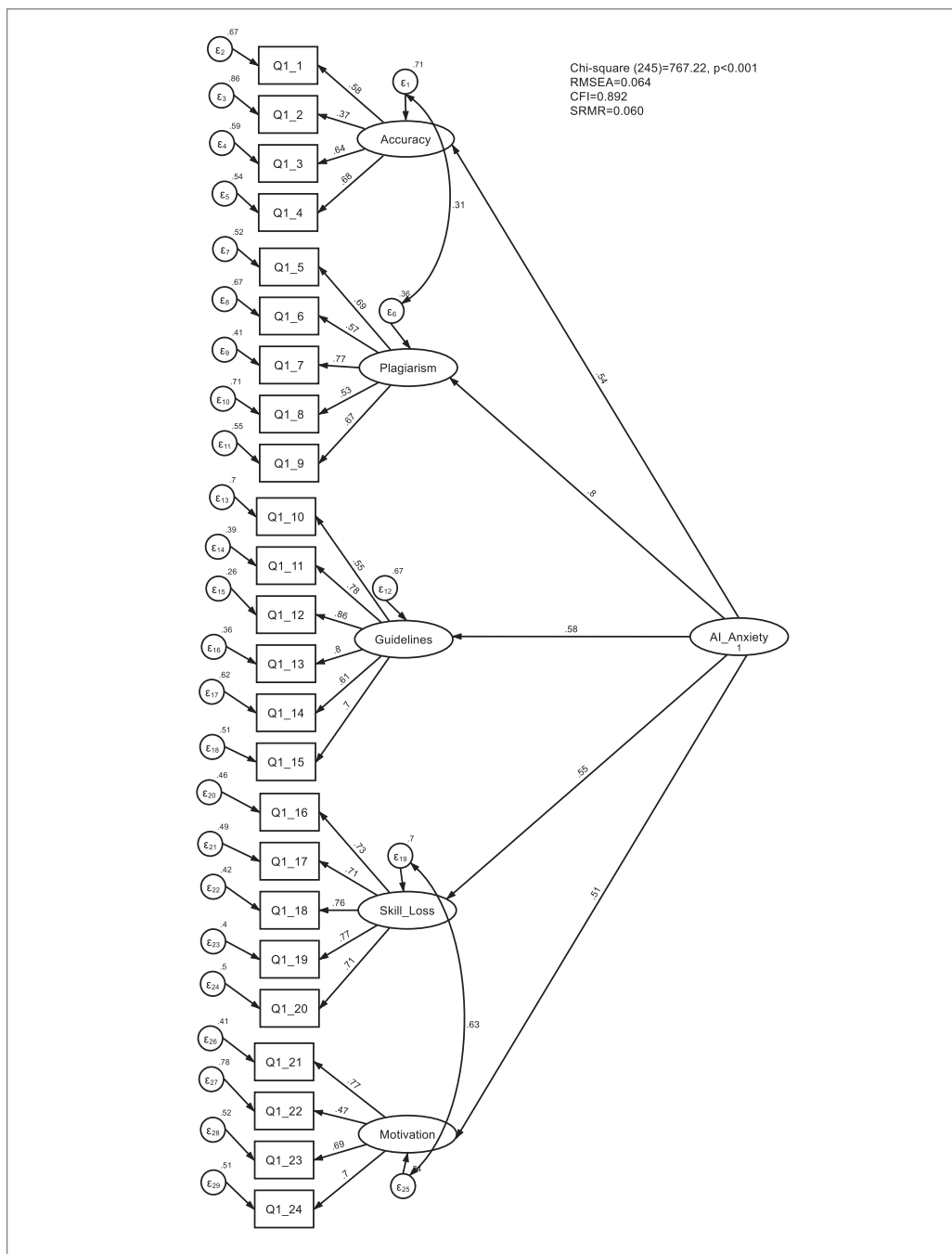


Рис. 1. Модель ИИ-тревожности
 Fig. 1: Artificial intelligence anxiety model

Примечание: Accuracy – «точность», Plagiarism – «плагиат», Guideline – «руководство», Skill_loss – «утрата навыков», Motivation – «мотивация», AI-Anxiety – «ИИ-тревожность».

Таблица 5

Результаты описательной статистики методики

Table 5

Descriptive Statistics for the Scale

Субшкалы ИИ-тревожности	М	SD
Руководство	2,79	1,06
Плагиаг	2,65	0,96
Мотивация	2,84	1,07
Утрата навыков	3,31	1,06
Точность	3,00	0,84

будущее» ($r = -,134; p = ,002$). Три измерения ИИ-тревожности положительно связаны со шкалой «тёмного будущего»: снижение мотивации ($r = ,167; p = ,000$), утрата навыков ($r = ,155; p = ,000$), сомнения в точности результатов ИИ ($r = ,156; p = ,000$).

Все субшкалы ИИ-тревожности положительно связаны с негативным отношением к ИИ в академической среде в целом («Озабоченность негативными последствиями использования ИИ») и отрицательно – с положительным отношением к ИИ («Утилитарно-позитивное отношение к ИИ»). Это позволяет говорить о высоком уровне внешней валидности, поскольку полученные с помощью адаптированного опросника показатели ИИ-тревожности усиливают скептицизм в отношении к ИИ как инструмента для учебной деятельности, поддерживают низкий уровень компетентности в работе с ИИ, неверие в возможности ИИ быть полезным ресурсом в обучении, опасения в том, что использование ИИ формирует лень и потерю интереса к учёбе, отнимает время и не приносит пользы, блокирует творческие навыки и личностное развитие, создаёт риски для будущей карьеры и образовательной справедливости. Напротив, снижает ИИ-тревожность активное и уверенное владение навыками работы с ИИ, убежденность в том, что ИИ расширяет возможности для творчества и личного развития студентов с разными когнитивными способностями (образовательная справедливость), повышает учебную продуктивность и мотивацию, по-

могает справляться с учебными нагрузками, будет полезен в будущей профессии, достоин доверия и в целом улучшает жизнь.

Результаты взаимосвязи ИИ-тревожности и личностных черт позволили выявить следующие факты. Опасения, связанные с академической честностью, выражены у людей с доброжелательностью ($r = ,083; p = ,058$) и добросовестностью ($r = ,085; p = ,052$). Напротив, добросовестность и устойчивость (в противовес нейротизму) снижает такой показатель ИИ-тревожности, как негативное влияние на учебную мотивацию ($r = -,153; p = ,000$ и $r = -,134; p = ,002$ соответственно). Кроме того, нейротизм связан с опасениями относительно утраты когнитивных навыков ($r = ,128; p = ,003$) и относительно точности данных ИИ ($r = ,158; p = ,000$).

Выявлено отсутствие корреляций с уровнем дохода и возрастом, что частично можно объяснить слабой вариативностью этих параметров в выборке. Не обнаружены и различия в ИИ-тревожности у студентов из разных населённых пунктов и разных курсов обучения. Уровень религиозности оказался положительно связан с опасениями относительно академической честности при использовании ИИ ($r = ,129; p = ,003$) и с негативным влиянием на учебную мотивацию ($r = ,085; p = ,052$). Девушки-студенты в большей степени, чем юноши, обеспокоены тем, как ИИ повлияет на академическую честность ($p = ,025$). Они испытывают более выраженные опасения относительно точности

и надёжности информации, предоставляемой ИИ ($p = ,007$), а также возможной утраты когнитивных навыков ($p = ,058$).

Описательная статистика субшкал приведена в *таблице 5*.

Обсуждение

Полученные результаты подтверждают, что ИИ-тревожность у студентов представляет собой сложное образование, структура которого в целом согласуется с данными, представленными в зарубежных исследованиях, однако обладает выраженной спецификой, обусловленной образовательным контекстом. Адаптированный и модифицированный психодиагностический инструмент обладает достаточным уровнем валидности и надёжности, и может применяться для исследований на российской студенческой выборке. Единственным ограничением психометрических данных следует считать пограничный коэффициент надёжности для шкалы «точность», что требует проверки в дальнейших исследованиях.

Особого внимания заслуживает тот факт, что наибольший вклад в общий уровень ИИ-тревожности вносит фактор, связанный с академической честностью (плагиат). Данный результат согласуется с выводами исследований, в которых подчёркивается роль нормативной неопределённости, серых зон в институциональных требованиях как источников тревоги при внедрении ИИ [43]. В образовательной среде ИИ оказывается включённым в систему оценки, санкций и ожиданий, вследствие чего тревожность формируется не только как реакция на технологическую неопределённость, но и как ответ на риск нарушения правил и последствий этого нарушения. Таким образом, ИИ-тревожность студентов носит в значительной степени институционально опосредованный характер.

Выделение фактора, связанного с отсутствием чётких правил использования ИИ, также подтверждает данные литературы о том, что неопределённость выступает одним

из ключевых психологических механизмов формирования тревоги [44]. В ряде работ показано, что недостаточная регламентация технологий усиливает чувство уязвимости [45] и снижает предсказуемость среды [46].

Фактор утраты когнитивных и учебных навыков, выявленный в исследовании, также находит подтверждение в литературе, где подчёркивается связь ИИ-тревожности с опасениями деградации профессиональных и интеллектуальных компетенций [47]. Данный аспект может быть рассмотрен относительно представлений о самоэффективности [48; 49; 50]: использование внешних инструментов, способных выполнять когнитивные функции, может восприниматься как угроза собственной компетентности, особенно в ситуациях, где ценится самостоятельное выполнение задач. Это позволяет интерпретировать ИИ-тревожность не только как страх перед технологией, но и как реакцию на возможное изменение структуры субъективного контроля над деятельностью. Подтверждение этому можно найти в зафиксированных связях ИИ-тревожности студентов с представлениями о будущем как управляемом и контролируемом. Здесь уместно провести параллели с исследованиями психологической «воодушевленности» и ощущения компетентности, влияющими на снижение тревоги, но сопровождающимися более критической рефлексией и настороженностью [25; 26].

Особую значимость имеет выделение мотивационного фактора, отсутствующего в оригинальной версии шкалы. Его появление подтверждает тезис о том, что ИИ-тревожность затрагивает не только когнитивные и нормативные аспекты, но и процессы саморегуляции деятельности. В исследовательской литературе отмечается, что взаимодействие с ИИ может изменять учебные стратегии, снижать усилия и влиять на поведенческие намерения и прежде всего на мотивацию [51; 52]. Полученные результаты конкретизируют эти положения, показывая, что тревожность может быть связана с про-

крастинацией, снижением вовлеченности и неуверенностью в собственных знаниях, то есть с изменением субъективной позиции учащегося в образовательном процессе.

В более широком контексте результаты исследования позволяют рассматривать ИИ-тревожность как феномен, связанный с переживанием будущего и трансформацией представлений о собственной роли в нем. Как отмечается в научной литературе, страхи, связанные с ИИ, затрагивают базовые психологические «буферы тревоги» – ощущение предсказуемости мира, значимости собственной деятельности и устойчивости идентичности [53–57]. В этом смысле ИИ-тревожность студентов может быть интерпретирована не только как реакция на конкретные образовательные изменения, но и как частный случай более широкой экзистенциальной тревоги, возникающей в условиях технологических трансформаций.

Факты о связи ИИ-тревожности с отношением к ИИ, в том числе готовностью пользоваться им для решения учебных заданий, согласуются с данными о технологической готовности студентов [19] и других социальных категорий [20; 21]. Можно утверждать, что низкий уровень технологической готовности работать с ИИ-инструментами будет повышать ИИ-тревожность студентов, и наоборот, компетентное обращение с ИИ приводит к ослаблению тревожности перед ним.

Полученные в исследовании данные подтверждают связь ИИ-тревоги с нейротизмом как личностной чертой [20]. В то же время и другие черты личности, особенно добросовестность, имеют значение для управления ИИ-тревожностью студентов. Это наблюдение требует дальнейших исследований в области профилактики ИИ-тревожности в академической среде, с обращением особого внимания на сверхответственных студентов со склонностью к перфекционизму.

На выборке российских студентов в целом подтверждаются гендерные различия в ИИ-тревожности [1; 26; 28; 30; 31]. Девушки в большей степени, чем юноши, склонны к

ИИ-тревоге, однако не по всем параметрам. Это означает, что девушки-студентки нуждаются в более выраженной профилактической и коррекционной работе по снижению ИИ-тревожности.

Отсутствие связи ИИ-тревожности с возрастом и уровнем дохода может объясняться относительной однородностью выборки, однако это является ограничением, которое должно быть преодолено в будущих исследованиях. Наряду с этим изучения требует решение вопроса о потенциальных различиях ИИ-тревожности студентов разных направлений обучения [4; 6]. Интерес представляет связь некоторых показателей ИИ-тревожности с уровнем религиозности, что в дальнейшем может рассматриваться в религиозном контексте использования технологий [1; 34].

Результаты исследования подтверждают, что ИИ-тревожность в образовательной среде формируется на пересечении когнитивных, институциональных и мотивационных факторов. В отличие от более общих моделей, акцентирующих внимание на страхах автоматизации и потери работы, в студенческой выборке доминирующую позицию занимают академические практики – нормы добросовестности, требования к самостоятельности и неопределённость регулятивной среды. Это позволяет рассматривать ИИ-тревожность как специфическую форму адаптации к изменениям образовательного пространства, в которой тревога выступает не только деструктивным, но и потенциально регулятивным механизмом, сигнализирующим о необходимости выработки новых норм взаимодействия с технологиями.

Проведённый анализ не только уточняет природу ИИ-тревожности, но и даёт возможность конкретизировать направления её профилактики и коррекции в образовательной среде. Эти направления охватывают институциональный, компетентностный, оценочный и индивидуально-психологический уровни.

Первое направление связано со снижением институциональной неопределённости.

Тревожность студентов формируется не только как реакция на саму технологию, но и как ответ на размытость правил её допустимого использования в учебной работе. В этой связи образовательным организациям целесообразно выстраивать не декларативные, а операциональные регламенты: фиксировать на общеуниверситетском и факультетском уровнях допустимые и недопустимые формы обращения к генеративному ИИ, заранее обозначать правила для конкретных видов заданий и оценочных процедур, а также вводить понятный порядок раскрытия факта использования ИИ в тех случаях, когда такое использование разрешено. Международная практика показывает, что именно прозрачность правил, их письменная фиксация применительно к каждому оценочному мероприятию и ясное разграничение между допустимой помощью и нарушением академической добросовестности выступают сегодня базовым механизмом снижения тревоги и профилактики конфликтов в учебном процессе.

Второе направление – формирование ИИ-грамотности как встроенного, а не факультативного элемента учебной адаптации. Учитывая зафиксированные в исследовании опасения относительно точности и надёжности ИИ, а также утраты учебных и когнитивных навыков, профилактические меры должны быть ориентированы на развитие у студентов навыков критической проверки ответа модели, поиска первоисточников, распознавания фактических ошибок и «галлюцинаций», понимания ограничений алгоритма и различения собственной интеллектуальной работы и машинной подсказки. Продуктивными здесь оказываются сравнительные учебные форматы, в которых студенты выполняют задание сначала самостоятельно, затем с использованием ИИ, а после этого сопоставляют результаты, обсуждают искажения, упрощения аргументации и пределы делегирования мыслительных операций машине. Именно такие практики сегодня составляют ядро университетских

программ ИИ-грамотности, где акцент делается на критическом, ответственном и этически осмысленном использовании генеративного ИИ.

Третье направление касается перестройки способов оценивания. Поскольку значимая часть ИИ-тревожности концентрируется вокруг риска санкций, плагиата и неясных границ допустимого, профилактика должна включать такую организацию заданий, при которой оценивается не только итоговый продукт, но и понимание, ход рассуждения и степень самостоятельности студента. Это может предполагать включение устных компонентов, кратких рефлексивных комментариев о том, как именно использовался ИИ, опору на материалы конкретного курса, обсуждения в аудитории, анализ визуальных и мультимодальных данных, а также задания, требующие собственной позиции и аргументации применительно к недавним или локально заданным кейсам. Зарубежные рекомендации показывают, что подобный переход от проверки «готового текста» к оценке понимания и процесса не только поддерживает академическую честность, но и снижает тревогу студентов, делая правила работы с ИИ более предсказуемыми и педагогически осмысленными.

Наконец, четвёртое направление – адресный, а не универсально-усредненный характер профилактических программ. Установленная в исследовании связь ИИ-тревожности с нейротизмом и добросовестностью, а также подтверждённые половые различия указывают на необходимость поддержки студентов, склонных к повышенной самокритичности, перфекционизму, сомнениям в допустимости даже разрешённых инструментов и более выраженной тревоге по поводу точности ИИ и академической честности. Практически это может реализовываться через вводные консультации по правилам использования ИИ, разбор типичных спорных ситуаций, наставничество со стороны старшекурсников, регулярные разъяснительные сессии с преподавателями, а также

периодическое обновление университетских руководств (*guidance*-документов) с учётом изменения технологий и учебных практик. Международные данные показывают, что наиболее продуктивными оказываются не разовые запреты, а сочетание понятных правил, кампаний по информированию студентов, диалога преподавателей и обучающихся и включения ИИ-грамотности в академические процессы.

Выводы

Полученные результаты позволяют заключить, что адаптированная и модифицированная версия опросника тревожности, связанной с развитием искусственного интеллекта, обладает удовлетворительными психометрическими характеристиками, а выявленная пятифакторная структура отражает специфику восприятия искусственного интеллекта студентами. Существенно, что ИИ-тревожность в образовательной среде формируется преимущественно вокруг конкретных академических практик – норм добросовестности, требований к самостоятельности и неопределённости регулятивной среды, что задаёт направление для дальнейших эмпирических и прикладных исследований.

Литература

1. Alkhalifah J.M., Bedaiwi A.M., Shaikh N., Seddiq W., Meo S.A. Existential anxiety about artificial intelligence (AI) – is it the end of humanity era or a new chapter in the human revolution: questionnaire-based observational study // *Frontiers in Psychiatry*. – 2024. – Vol. 9, no. 15. – Article no. 1368122. – DOI: 10.3389/fpsy.2024.1368122.
2. Rashid A.B., Kausik, M.A.K. AI Revolutionizing Industries Worldwide: A Comprehensive Overview of Its Diverse Applications // *Hybrid Advances*. – 2024. – Vol. 7. – Article no. 100277. – DOI: 10.1016/j.hybadv.2024.100277.
3. Fry A.D.J. Being Human in a Post-Human World: Actor-Network Theory for Anxiety-Induced Meaning-Making in the Existential Shadow of AI // *Implicit Religion*. – 2025. – Vol. 26, no. 2–3. – P. 231–257. – DOI: 10.1558/imre.30727.
4. Rizkina A.T., Hidajat H.G., Farida I.A. Job-Related Anxiety in the Age of Artificial Intelligence: A Systematic Review of Workplace Dynamics // *Formosa Journal of Multidisciplinary Research*. – 2025. – Vol. 4, no. 8. – P. 3965–3976. – DOI: 10.55927/fjmr.v4i8.416.
5. Hatos A. Anxiety in the Age of AI: Constructing a Tool to Assess Public Perceptions // *BRAIN. Broad Research In Artificial Intelligence And Neuroscience*. – 2025. – Vol. 16, no. 1. – P. 415–426. – DOI: 10.70594/brain/16.S1/32.
6. Dong M., Conway J.R., Bonnefon J.F., Shariff A., Rahwan I. Fears about artificial intelligence across 20 countries and six domains of application // *American Psychologist*. – 2026. – Vol. 81, no. 1. – P. 53–67. – DOI: 10.1037/amp0001454.
7. Suseno Y., Chang C., Hudik M., Fang E.S. Beliefs, Anxiety and Change Readiness for Artificial Intelligence Adoption among Human Resource Managers: The Moderating Role of High-Performance Work Systems. In A. Malik, & P. Budhwar (Eds.) // *Artificial Intelligence and International HRM*. – 2023. – P. 144–171. – DOI: 10.4324/9781003377085-6.
8. Song K., Guo M., Ye L., Liu Y., Liu S. Driverless metros are coming, but what about the drivers? A study on AI-related anxiety and safety performance // *Safety Science*. – 2024. – Vol. 175. – Article no. 106487. – DOI: 10.1016/j.ssci.2024.106487.
9. Goethals F., Ziegelmayer J.L. Anxiety buffers and the threat of extreme automation: a terror management theory perspective // *Information Technology & People*. – 2022. – Vol. 35, no. 1. – P. 96–118. – DOI: 10.1108/ITP-06-2019-0304.
10. Rony M.K.K., Parvin M.R., Wahiduzzaman M., Debnath M., Bala S.D., Kayesh I. “I Wonder if my Years of Training and Expertise Will be Devalued by Machines”: Concerns About the Replacement of Medical Professionals by Artificial Intelligence // *SAGE Open Nurs*. – 2024. – Vol. 7, no. 10. – DOI: 10.1177/23779608241245220.
11. Sharma V., Deb S., Mahajan Y., Ghosal A., Kapse M. Psychological impacts of AI-induced job displacement among Indian IT professionals: a Delphi-validated thematic analysis // *International Journal of Qualitative Studies on Health and Well-being*. – 2025. – Vol. 20, no. 1. – Article no. 2556445. – DOI: 10.1080/17482631.2025.2556445.
12. Wang Y.Y., Wang Y.S. Development and validation of an artificial intelligence anxiety scale: an

- initial application in predicting motivated learning behavior // *Interactive Learning Environments*. – 2019. – Vol. 30, no. 4. – P. 619–634. – DOI: 10.1080/10494820.2019.1674887.
13. Zhan E.S., Molina M.D., Rheu M., Peng W. What is there to fear? Understanding multi-dimensional fear of ai from a technological affordance perspective // *International Journal of Human-Computer Interaction*. – 2023. – DOI: 10.1080/10447318.2023.2261731.
 14. Wang Y.M., Wei C.L., Lin H.H., Wang S.C., Wang Y.S. What drives students' AI learning behavior: a perspective of AI anxiety // *Interactive Learning Environments*. – 2024. – Vol. 32, no. 6. – P. 2584–2600. – DOI: 10.1080/10494820.2022.2153147.
 15. Wen F., Li, Y., Zhou Y., An X., Zou Q. A Study on the Relationship between AI Anxiety and AI behavioral intention of secondary school students learning English as a foreign language // *Journal of Educational Technology Development and Exchange*. – 2024. – Vol. 17, no. 1. – P. 130–154. – DOI: 10.18785/jetde.1701.07.
 16. Liu X., Liu Y. Developing and Validating a Scale of Artificial Intelligence Anxiety Among Chinese EFL Teachers // *European Journal of Education*. – 2025. – Vol. 60, iss. 1. – DOI: 10.1111/ejed.12902.
 17. Cugurullo F., Acheampong R.A. Fear of AI: an inquiry into the adoption of autonomous cars in spite of fear, and a theoretical framework for the study of artificial intelligence technology acceptance // *AI & Society*. – 2024. – Vol. 39. – P. 1569–1584. – DOI: 10.1007/s00146-022-01598-6.
 18. Kitchens M.B., Meier B.P. The fearful mind of artificial intelligence: fear and perceived existential threat of artificial intelligence as a function of its cognitive and emotional capabilities // *Journal of Social Psychology*. – 2025. – Vol. 9. – P. 1–14. – DOI: 10.1080/00224545.2025.2503006.
 19. Lemay D., Basnet R.B., Doleck T. Fearing the Robot Apocalypse: Correlates of AI Anxiety // *International Journal of Learning Analytics and Artificial Intelligence for Education*. – 2020. – Vol. 2, no. 2. – P. 24–33. – DOI: 10.3991/ijai.v2i2.16759.
 20. Kaya F., Aydin F., Schepman A., Rodway P., Yetişensoy O., Demir Kaya M. The Roles of Personality Traits, AI Anxiety, and Demographic Factors in Attitudes toward Artificial Intelligence // *International Journal of Human-Computer Interaction*. – 2024. – No. 40. – P. 497–514. – DOI: 10.1080/10447318.2022.2151730.
 21. Schiavo G., Businaro S., Zancanaro M. Comprehension, Apprehension, and Acceptance: Understanding the Influence of Literacy and Anxiety on Acceptance of Artificial Intelligence // *Technology and Society*. – 2024. – Vol. 77. – DOI: 10.1016/j.techsoc.2024.102537.
 22. Zhong G. The influence mechanism of ai technology learning anxiety in the human-computer collaboration context: the moderating effect of uncertainty avoidance and the mediating effect of self-efficacy // *The EUrASEANs: Journal on Global Socio-Economic Dynamics* – 2023. – No. 6 (43). – DOI: 10.35678/2539-5645.6(43).2023.170-180.
 23. Chang P.C., Zhang W., Cai Q., Guo H. Does AI-Driven Technostress Promote or Hinder Employees' Artificial Intelligence Adoption Intention? A Moderated Mediation Model of Affective Reactions and Technical Self-Efficacy // *Psychology Research and Behavior Management*. – 2024. – Vol. 7, no. 17. – P. 413–427. – DOI: 10.2147/PRBM.S441444.
 24. Tsao J. Trajectories of AI policy in higher education: Interpretations, discourses, and enactments of students and teachers // *Computers and Education: Artificial Intelligence*. – 2025. – DOI: 10.1016/j.caeai.2025.100496.
 25. Onan G., Dalmış A.B. Yapay Zeka Kaygısının Gölgesinde: Psikolojik Güçlendirmenin İşe Adanmışlık Ve İşte Kendini Yetiştirme Üzerindeki Etkisi // *Business & Management Studies: An International Journal*. – 2025. – Vol. 13, no. 1. – P. 300–321. – DOI: 10.15295/bmij.v13i1.2508.
 26. Güdük O., Vural A., Dişiaçık G. Investigating The Effect of Perceived Empowerment on Artificial Intelligence Anxiety Levels in Healthcare Workers" // *Çalışma Ve Toplum*. – 2025. – Vol. 1, iss. 84. – P. 285–310. – DOI: 10.54752/ct.1624015.
 27. Beder Z.A., Donmez-Turan A. The Critical Role of Uncertainty Intolerance on the Relationship Between Spiritual Intelligence and Artificial Intelligence Anxiety // *International Journal of Business Ecosystem & Strategy*. – 2025. – Vol. 7, no. 1. – P. 82–105. – DOI: 10.36096/ijbes.v7i1.722.
 28. Russo C., Romano L., Clemente D., Iacovone L., Gladwin T.E., Panno A. Gender differences in artificial intelligence: the role of artificial intelligence anxiety // *Frontiers in Psychology*. – 2025. – Vol. 16. – Article no. 1559457. – DOI: 10.3389/fpsyg.2025.1559457.

29. Terzi R. An adaptation of artificial intelligence anxiety scale into turkish: reliability and validity study // *International Online Journal of Education and Teaching (IOJET)*. – 2020. – Vol. 7, no. 4. – P. 1501–1515. – DOI: 10.1177/10519815241290648.
30. Falebita O. Evaluating Artificial Intelligence Anxiety Among Pre-Service Teachers in University Teacher Education Programs // *Journal of Mathematics Instruction, Social Research and Opinion*. – 2024. – Vol. 4, no. 1. – P. 1–16. – DOI: 10.58421/misro.v4i1.309.
31. Chen S., Zhao Y. Explore the AI Anxiety Attitude of E-Commerce Employees // *International Journal of Human-Computer Interaction*. – 2024. – Vol. 41, no. 13. – P. 8438–8446. – DOI: 10.1080/10447318.2024.2409470.
32. Gerlich M. Public Anxieties About AI: Implications for Corporate Strategy and Societal Impact // *Administrative Science*. – 2024. – Vol. 14. – 288 p. – DOI: 10.3390/admsci14110288.
33. Uçar M., Çapuk H., Yiğit M.F. The relationship between artificial intelligence anxiety and unemployment anxiety among university students // *Work*. – 2025. – Vol. 80, no. 2. – P. 701–710. – DOI: 10.1177/10519815241290648.
34. Arvai N., Katonai G., Mesko B. Health Care Professionals' Concerns About Medical AI and Psychological Barriers and Strategies for Successful Implementation: Scoping Review // *Journal of Medical Internet Research*. – 2025. – Vol. 23, no. 27. – Article no. e66986. – DOI: 10.2196/66986.
35. Alshaibani M.H., Al-Rahmi W.M., Tawafak R.M. Psychometric validation of the artificial intelligence anxiety scale: A confirmatory factor analysis for academic research // *Contemporary Educational Technology*. – 2025. – Vol. 17, no. 4. – DOI: 10.30935/cedtech/17309.
36. Сергеева А.С., Кириллов Б.А., Джумагулова А.Ф. Перевод и адаптация краткого пятифакторного опросника личности (TIPI-RU): оценка конвергентной валидности, внутренней согласованности и тест ретестовой надёжности // *Экспериментальная психология*. – 2016. – Т. 9, № 3. – С. 138–154. – DOI: 10.17759/exppsy.2016090311.
37. Zaleski Z., Sobol-Kwapinska M., Przepiorka A., Meisner M. Development and validation of the Dark Future scale // *Time & Society*. – 2019. – Vol. 28, no. 1. – P. 107–123. – DOI: 10.1177/0961463X16678257.
38. Нестик Т.А., Журавлев А.А. Психология глобальных рисков. – Москва: Институт психологии РАН, 2018. – 402 с. – ISBN: 978-5-9270-0385-3.
39. Нестик Т.А. Коллективный образ будущего: социально-психологический анализ. – Москва: Институт психологии РАН, 2025. – 657 с. – DOI: 10.38098/soc_25_0493.
40. Hair J.F., Black W.C., Babin B.J. and Anderson R.E. *Multivariate Data Analysis*. Prentice Hall, Upper Saddle River, United States. – 2010. – 834 p. – ISBN: 978-1-4737-5654-0.
41. Acock A.C. *Discovering Structural Equation Modeling Using Stata*. Stata Press, College Station, United States. – 2013. – 306 p. – ISBN: 978-1-59718-133-4.
42. Byrne B.M. *Structural Equation Modeling with Amos*. Routledge, New York, United States. – 2016. – 396 p. – ISBN: 978-0-8058-6372-7.
43. Tao M., Li X., Alam F., Yan Y., Chau T. Unveiling the Impact of AI Technological Anxiety on the Marketers' Intention to Adopt Generative AI // *Journal of Global Information Management*. – 2025. – Vol. 33, no. 1. – P. 1–22. – DOI: 10.4018/JGIM.372177.
44. Gu Y., Gu S., Lei Y., Li H. From Uncertainty to Anxiety: How Uncertainty Fuels Anxiety in a Process Mediated by Intolerance of Uncertainty // *Neural Plasticity*. – 2020. – Nov 22. – DOI: 10.1155/2020/8866386.
45. Teo S.A. Artificial intelligence, human vulnerability and multi-level resilience // *Computer Law & Security Review*. – 2025. – Vol. 57. – Article no. 106134. – DOI: 10.1016/j.clsr.2025.106134.
46. Terenzio F. Systemic Vulnerability: From AI Systems to Environmental Systems // *Topoi*. – 2025. – DOI: 10.1007/s11245-025-10285-2.
47. Klein C.R., Klein R. The extended hollowed mind: why foundational knowledge is indispensable in the age of AI // *Frontiers in Artificial Intelligence*. – 2025. – Vol. 11, no. 8. – DOI: 10.3389/frai.2025.1719019.
48. Varol B. Artificial Intelligence Anxiety in Nursing Students: The Impact of Self-efficacy // *Computers, Informatics, Nursing*. – 2025. – Vol. 1, no. 43 (6). – Article no. e01250. – DOI: 10.1097/CIN.0000000000001250.
49. Morales-García W.C., Sairitupa-Sanchez L.Z., Flores-Paredes A., Pascual-Mariño J., Morales-García M. Influence of Self-Efficacy in the Use of Artificial Intelligence (AI) and Anxiety Toward AI Use on AI Dependence Among Peruvian Uni-

- versity Students. Data and Metadata [Internet]. 2025 – Jan. 12 [cited 2026 Apr. 3]; 4:210. – DOI: 10.56294/dm2025210.
50. Albino M.G., Albino F.S., Asio J.M.R., Gadia E.D. Influence of AI anxiety on AI self-efficacy among college students // *International Journal of Technology in Education (IJTE)*. – 2025. – Vol. 8, no. 2. – P. 557–573. – DOI: 10.46328/ijte.1109.
 51. Fan Y., Tang L., Le H., Shen K., Tan S., Zhao Y. et al. Beware of Metacognitive Laziness: Effects of Generative Artificial Intelligence on Learning Motivation, Processes, and Performance // *British Journal of Educational Technology*. – 2024. – Vol. 56. – P. 489–530. – DOI: 10.1111/bjet.13544.
 52. Kinskofer F., Tulis M. Motivational and appraisal factors shaping generative AI use and intention in Austrian higher education students and teachers // *Frontiers in Education*. – 2025. – Vol. 10. – DOI: 10.3389/educ.2025.1677827.
 53. Zhou J., Lu Y., Chen Q. GAI identity threat: When and why do individuals feel threatened? // *Information & Management*. – 2025. – Vol. 62, iss. 2. – Article no. 104093. – DOI: 10.1016/j.im.2024.104093.
 54. Zhou Q., Yang L., Tang Y., Yang J., Zhou W., Guan W., Yan L., Liu Y. The mediation of trust on artificial intelligence anxiety and continuous adoption of artificial intelligence technology among primacy nurses: a cross-sectional study // *BMC Nursing*. – 2025. – Vol. 1, iss. 24 (1). – Article no. 724. – DOI: 10.1186/s12912-025-03406-0.
 55. Andreescu F. Existence hacked: meaning, freedom, death, and intimacy in the age of AI // *AI & Society*. – 2024. – Vol. 40, no. 4. – P. 2881–2893. – DOI: 10.1007/s00146-024-02052-5.
 56. Jussupow E., Spohrer K., Heinzl A. Identity Threats as a Reason for Resistance to Artificial Intelligence: Survey Study With Medical Students and Professionals // *JMIR Formative Research*. – 2022. – Vol. 23, no. 6 (3). – Article no. e28750. – DOI: 10.2196/28750.
 57. Kuchmaner C.A. Mitigating Identity Threats Posed by Artificial Intelligence: The Role of Perceived AI Relatedness // *Journal of Consumer Behavior*. – 2026. – Vol. 25, no. 1. – P. 176–190. – DOI: 10.1002/cb.70066.

Статья поступила в редакцию 20.04.2026

Принята к публикации 29.05.2026

References

1. Alkhalifah, J.M., Bedaiwi, A.M., Shaikh, N., Seddiq, W., Meo, S.A. (2024). Existential Anxiety about Artificial Intelligence (AI) - Is It the End of Humanity Era or a New Chapter in the Human Revolution: Questionnaire-Based Observational Study. *Frontiers in Psychiatry*. Vol. 9, no. 15, article no. 1368122, doi: 10.3389/fpsy.2024.1368122.
2. Rashid, A.B., Kausik, M.A.K. (2024). AI Revolutionizing Industries Worldwide: A Comprehensive Overview of Its Diverse Applications. *Hybrid Advances*, Vol. 7, article no. 100277, doi: 10.1016/j.hybadv.2024.100277.
3. Fry, A.D.J. (2025). Being Human in a Post-Human World: Actor-Network Theory for Anxiety-Induced Meaning-Making in the Existential Shadow of AI. *Implicit Religion*. Vol. 26, no. 2-3, pp. 231-257, doi: 10.1558/imre.30727.
4. Rizkina, A.T., Hidajat, H.G., Farida, I.A. (2025). Job-Related Anxiety in the Age of Artificial Intelligence: A Systematic Review of Workplace Dynamics. *Formosa Journal of Multidisciplinary Research*. Vol. 4, no. 8. pp. 3965-3976, doi: 10.55927/fjmr.v4i8.416.
5. Hatos, A. (2025). Anxiety in the Age of AI: Constructing a Tool to Assess Public Perceptions. *BRAIN. Broad Research In Artificial Intelligence And Neuroscience*. Vol. 16, pp. 415-426, doi: 10.70594/brain/16.S1/32.
6. Dong, M., Conway, J.R., Bonnefon, J.F., Shariff, A., Rahwan, I. (2026). Fears about Artificial Intelligence across 20 Countries and Six Domains of Application. *American Psychologist*. Vol. 81, no. 1, pp. 53-67, doi: 10.1037/amp0001454.
7. Suseno, Y., Chang, C., Hudik, M., Fang, E.S. (2023). Beliefs, Anxiety and Change Readiness for Artificial Intelligence Adoption among Human Resource Managers: The Moderating Role of

- High-Performance Work Systems. In: A. Malik, P. Budhwar (Eds.). *Artificial Intelligence and International HRM*. Pp. 144-171, doi: 10.4324/9781003377085-6.
8. Song, K., Guo, M., Ye, L., Liu, Y., Liu, S. (2024). Driverless Metros Are Coming, but What about the Drivers? A Study on AI-Related Anxiety and Safety Performance. *Safety Science*. Vol. 175, article no. 106487, doi: 10.1016/j.ssci.2024.106487.
 9. Goethals, F., Ziegelmayer, J.L. (2022). Anxiety Buffers and the Threat of Extreme Automation: A Terror Management Theory Perspective. *Information Technology & People*. Vol. 35, no. 1, pp. 96-118, doi: 10.1108/ITP-06-2019-0304.
 10. Rony, M.K.K., Parvin, M.R., Wahiduzzaman, M., Debnath, M., Bala, S.D., Kayesh, I. (2024). "I Wonder if my Years of Training and Expertise Will be Devalued by Machines": Concerns About the Replacement of Medical Professionals by Artificial Intelligence. *SAGE Open Nurs*. Vol. 7, no. 10, doi: 10.1177/23779608241245220.
 11. Sharma, V., Deb S., Mahajan, Y., Ghosal, A., Kapse, M. (2025). Psychological Impacts of AI-Induced Job Displacement among Indian IT Professionals: A Delphi-Validated Thematic Analysis. *Int J Qual Stud Health Well-being*. Vol. 20, no. 1, article no. 2556445, doi: 10.1080/17482631.2025.2556445.
 12. Wang, Y.Y., Wang, Y.S. (2019). Development and Validation of an Artificial Intelligence Anxiety Scale: An Initial Application in Predicting Motivated Learning Behavior. *Interactive Learning Environments*. Vol. 30, no. 4, pp. 619-634, doi: 10.1080/10494820.2019.1674887.
 13. Zhan, E.S., Molina, M.D., Rheu, M., Peng, W. (2023). What is There to Fear? Understanding Multi-Dimensional Fear of Ai from a Technological Affordance Perspective. *International Journal of Human-Computer Interaction*. Doi: 10.1080/10447318.2023.2261731.
 14. Wang, Y.M., Wei, C.L., Lin, H.H., Wang, S.C., Wang, Y.S. (2024). What Drives Students' AI Learning Behavior: A Perspective of AI Anxiety. *Interactive Learning Environments*. Vol. 32, no. 6, pp. 2584-2600, doi: 10.1080/10494820.2022.2153147.
 15. Wen, F., Li, Y., Zhou, Y., An, X., Zou, Q. (2024). A Study on the Relationship between AI Anxiety and AI Behavioral Intention of Secondary School Students Learning English as a Foreign Language. *Journal of Educational Technology Development and Exchange*. Vol. 17, no. 1, pp. 130-154, doi: 10.18785/jetde.1701.07.
 16. Liu, X., Liu, Y. (2025). Developing and Validating a Scale of Artificial Intelligence Anxiety Among Chinese EFL Teachers. *European Journal of Education*. Vol. 60, iss. 1, doi: 10.1111/ejed.12902.
 17. Cugurullo, F., Acheampong, R.A. (2024). Fear of AI: An Inquiry into the Adoption of Autonomous Cars in Spite of Fear, and a Theoretical Framework for the Study of Artificial Intelligence Technology Acceptance. *AI & Society*. Vol. 39, pp. 1569-1584, doi: 10.1007/s00146-022-01598-6.
 18. Kitchens, M.B., Meier, B.P. (2025). The Fearful Mind of Artificial Intelligence: Fear and Perceived Existential Threat of Artificial Intelligence as a Function of Its Cognitive and Emotional Capabilities. *Journal of Social Psychology*. Vol. 9, pp. 1-14, doi: 10.1080/00224545.2025.2503006.
 19. Lemay, D., Basnet, R.B., Doleck, T. (2020). Fearing the Robot Apocalypse: Correlates of AI Anxiety. *International Journal of Learning Analytics and Artificial Intelligence for Education*. Vol. 2, no. 2, pp. 24-33, doi: 10.3991/ijai.v2i2.16759.
 20. Kaya, F., Aydin, F., Schepman, A., Rodway, P., Yetişensoy, O., Demir Kaya, M. (2024). The Roles of Personality Traits, AI Anxiety, and Demographic Factors in Attitudes toward Artificial Intelligence. *International Journal of Human-Computer Interaction*. Vol. 40, pp. 497-514, doi: 10.1080/10447318.2022.2151730.
 21. Schiavo, G., Businaro, S., Zancanaro, M. (2024). Comprehension, Apprehension, and Acceptance: Understanding the Influence of Literacy and Anxiety on Acceptance of Artificial Intelligence. *Technology and Society*. Vol. 77, doi: 10.1016/j.techsoc.2024.102537.

22. Zhong, G. (2023). The Influence Mechanism of Ai Technology Learning Anxiety in the Human-Computer Collaboration Context: The Moderating Effect of Uncertainty Avoidance and the Mediating Effect of Self-Efficacy. *The EURASEANs: Journal on Global Socio-Economic Dynamics*. No. 6 (43), doi: 10.35678/2539-5645.6(43).2023.170-180.
23. Chang, P.C., Zhang, W., Cai, Q., Guo, H. (2024). Does AI-Driven Technostress Promote or Hinder Employees' Artificial Intelligence Adoption Intention? A Moderated Mediation Model of Affective Reactions and Technical Self-Efficacy. *Psychology Research and Behavior Management*. Vol. 7, no. 17, doi: 10.2147/PRBM.S441444.
24. Tsao, J. (2025). Trajectories of AI Policy in Higher Education: Interpretations, Discourses, and Enactments of Students and Teachers. *Computers and Education: Artificial Intelligence*. Doi: 10.1016/j.caeai.2025.100496.
25. Onan, G., Dalmış, A.B. (2025). Yapay Zekâ Kaygısının Gölgesinde: Psikolojik Güçlendirmenin İşe Adanmışlık Ve İşte Kendini Yetiştirme Üzerindeki Etkisi. *Business & Management Studies: An International Journal*. Vol. 13, no. 1, pp. 300-321, doi: 10.15295/bmij.v13i1.2508.
26. Güdük, O., Vural, A., Dişiaçık, G. (2025). Investigating The Effect of Perceived Empowerment on Artificial Intelligence Anxiety Levels in Healthcare Workers. *Çalışma Ve Toplum*. Vol. 1, iss. 84, pp. 285-310, doi: 10.54752/ct.1624015.
27. Beder, Z. A., Donmez-Turan, A. (2025). The Critical Role of Uncertainty Intolerance on the Relationship Between Spiritual Intelligence and Artificial Intelligence Anxiety. *International Journal of Business Ecosystem & Strategy* (2687-2293). Vol. 7, no. 1, pp. 82-105. doi: 10.36096/ijbes.v7i1.722.
28. Russo, C., Romano, L., Clemente, D., Iacovone L., Gladwin, T.E., Panno, A. (2025). Gender Differences in Artificial Intelligence: The Role of Artificial Intelligence Anxiety. *Frontiers in Psychology*. Vol. 6, no. 16, doi: 10.3389/fpsyg.2025.1559457.
29. Terzi, R. (2020). An Adaptation of Artificial Intelligence Anxiety Scale into Turkish: Reliability and Validity Study. *International Online Journal of Education and Teaching (IOJET)*. Vol. 7, no. 4, pp. 1501-1515, doi: 10.1177/10519815241290648.
30. Falebita, O. (2024). Evaluating Artificial Intelligence Anxiety Among Pre-Service Teachers in University Teacher Education Programs. *Journal of Mathematics Instruction, Social Research and Opinion*. Vol. 4, no. 1, pp. 1-16, doi: 10.58421/misro.v4i1.309.
31. Chen, S., Zhao, Y. (2024). Explore the AI Anxiety Attitude of E-Commerce Employees. *International Journal of Human-Computer Interaction*. Vol. 41, no. 13, pp. 8438-8446, doi: 10.1080/10447318.2024.2409470.
32. Gerlich, M. (2024). Public Anxieties About AI: Implications for Corporate Strategy and Societal Impact. *Administrative Science*. Vol. 14, no. 288, doi: 10.3390/admsci14110288.
33. Uçar, M., Çapuk, H., Yigit, M.F. (2025). The Relationship Between Artificial Intelligence Anxiety and Unemployment Anxiety among University Students. *Work*. Vol. 80, no. 2, pp. 701-710, doi: 10.1177/10519815241290648.
34. Arvai, N., Katonai, G., Mesko, B. (2025). Health Care Professionals' Concerns About Medical AI and Psychological Barriers and Strategies for Successful Implementation: Scoping Review. *Journal of Medical Internet Research*. Vol. 23, no. 27, article no. e66986, doi: 10.2196/66986.
35. Alshaibani, M.H., Al-Rahmi, W.M., Tawafak, R.M. (2025). Psychometric Validation of the Artificial Intelligence Anxiety Scale: A Confirmatory Factor Analysis for Academic Research. *Contemporary Educational Technology*. Vol. 17, no. 4, doi: 10.30935/cedtech/17309.
36. Sergeeva, A.S., Kirillov, B.A., Dzhumagulova, A.F. (2016). Translation and Adaptation of Short Five Factor Personality Questionnaire (Tipi-Ru): Convergent Validity, Internal Consistency and

- Test-Retest Reliability Evaluation. *Eksperimental' naya psikhologiya = Experimental psychology*. Vol. 9, no. 3. pp. 138-154, doi: 10.17759/exppsy.2016090311 (In Russ., abstract in Eng.).
37. Zaleski, Z., Sobol-Kwapinska, M., Przepiorka, A., Meisner, M. (2019). Development and Validation of the Dark Future Scale. *Time & Society*, Vol. 28, no. 1, pp. 107-123, doi: 10.1177/0961463X16678257.
 38. Nestik, T.A., Zhuravlev, A.L. (2018). *Psychology of Global Risks* [Psychology of Global Risks]. Moscow: Institute of Psychology of the Russian Academy of Sciences. 402 p. ISBN: 978-5-9270-0385-3. (In Russ.).
 39. Nestik T.A. (2025). *Kollektivnyi obraz budushchego: sotsial'no-psikhologicheskii analiz* [The Collective Image of the Future: A Socio-Psychological Analysis]. Moscow: Institut psikhologii RAN, 657 p., doi: 10.38098/soc_25_0493 (In Russ.).
 40. Hair, J.F., Black, W.C., Babin, B.J., Anderson, R.E. (2010). *Multivariate Data Analysis*. Prentice Hall, Upper Saddle River, United States, 834 p. ISBN: 978-1-4737-5654-0.
 41. Acock, A.C. (2013). *Discovering Structural Equation Modeling Using Stata*. Stata Press, College Station, United States, 306 p. ISBN: 978-1-59718-133-4.
 42. Byrne, B.M. (2016). *Structural Equation Modeling with Amos*. Routledge, New York, United States, 396 p. ISBN 978-0-8058-6372-7.
 43. Tao, M., Li, X., Alam, F., Yan, Y., Chau, T. (2025). Unveiling the Impact of AI Technological Anxiety on the Marketers' Intention to Adopt Generative AI. *Journal of Global Information Management*. Vol. 33, no. 1. pp. 1-22. doi: 10.4018/JGIM.372177.
 44. Gu, Y, Gu, S, Lei, Y, Li, H. (2020). From Uncertainty to Anxiety: How Uncertainty Fuels Anxiety in a Process Mediated by Intolerance of Uncertainty. *Neural Plasticity*. Vol. 22, doi: 10.1155/2020/8866386.
 45. Teo, S.A. (2025). Artificial Intelligence, Human Vulnerability and Multi-Level Resilience. *Computer Law & Security Review*. Vol. 57, article no. 106134, doi: 10.1016/j.clsr.2025.106134.
 46. Terenzio, F. (2025). Systemic Vulnerability: From AI Systems to Environmental Systems. *Topoi*. doi: 10.1007/s11245-025-10285-2.
 47. Klein, C.R., Klein, R. (2025). The Extended Hollowed Mind: Why Foundational Knowledge Is Indispensable in the Age of AI. *Frontiers in Artificial Intelligence*. Vol. 11, no. 8, doi: 10.3389/frai.2025.1719019.
 48. Varol, B. (2025). Artificial Intelligence Anxiety in Nursing Students: The Impact of Self-efficacy. *Computers, Informatics, Nursing*. Vol. 1, no. 43 (6), article no. e01250, doi: 10.1097/CIN.0000000000001250.
 49. Morales-García, W.C., Sairitupa-Sanchez, L.Z., Flores-Paredes, A., Pascual-Mariño, J., Morales-García, M. (2026). *Influence of Self-Efficacy in the Use of Artificial Intelligence (AI) and Anxiety Toward AI Use on AI Dependence Among Peruvian University Students*. Data and Metadata [Internet]. 2025 Jan. 12 [cited 2026 Apr. 3]; 4:210, doi: 10.56294/dm2025210.
 50. Albino, M.G., Albino, F.S., Asio, J.M.R., Gadia, E.D. (2025). Influence of AI Anxiety on AI Self-Efficacy among College Students. *International Journal of Technology in Education (IJTE)*. Vol. 8, no. 2, pp. 557-573, doi: 10.46328/ijte.1109.
 51. Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y. et al. (2024). Beware of Metacognitive Laziness: Effects of Generative Artificial Intelligence on Learning Motivation, Processes, and Performance. *British Journal of Educational Technology*. Vol. 56, pp. 489-530, doi: 10.1111/bjet.13544.
 52. Kinskofer, F., Tulis, M. (2025) Motivational and Appraisal Factors Shaping Generative AI Use and Intention in Austrian Higher Education Students and Teachers. *Frontiers in Education*. Vol. 10, doi: 10.3389/educ.2025.1677827.

53. Zhou, J., Lu, Y., Chen, Q. (2025) GAI Identity Threat: When and Why Do Individuals Feel Threatened? *Information & Management*. Vol. 62, iss. 2, doi: 10.1016/j.im.2024.104093.
54. Zhou, Q., Yang, L., Tang, Y., Yang, J., Zhou, W., Guan, W., Yan, L., Liu Y. (2025). The Mediation of Trust on Artificial Intelligence Anxiety and Continuous Adoption of Artificial Intelligence Technology among Primacy Nurses: A Cross-Sectional Study. *BMC Nursing*. Vol. 1, iss. 24 (1), article no. 724, doi: 10.1186/s12912-025-03406-0.
55. Andreescu, F. (2024). Existence Hacked: Meaning, Freedom, Death, and Intimacy in the Age of AI. *Ai & Society*. Vol. 40 (4), pp. 2881-2893, doi: 10.1007/s00146-024-02052-5.
56. Jussupow, E., Spohrer., K, Heinzl., A. (2022). Identity Threats as a Reason for Resistance to Artificial Intelligence: Survey Study With Medical Students and Professionals. *JMIR Formative Research*. Vol. 23, no. 6 (3), doi: 10.2196/28750.
57. Kuchmaner, C.A. (2026). Mitigating Identity Threats Posed by Artificial Intelligence: The Role of Perceived AI Relatedness. *Journal of Consumer Behavior*. Vol. 25, no. 1, pp. 176-190, doi: 10.1002/cb.70066.

Статья поступила в редакцию 20.04.2026

Принята к публикации 29.05.2026

Опросник ИИ-тревожности (версия для студентов, 24 утверждения)

Инструкция: Оцените степень согласия с каждым утверждением по шкале от 1 до 5: 1 – полностью не согласен(а), 5 – полностью согласен(а).

1. Я беспокоюсь о точности ответов, которые предоставляет ИИ при выполнении учебных заданий.
2. Мне бывает трудно понять и интерпретировать результаты, сгенерированные ИИ.
3. Я опасаясь, что ИИ может давать информацию, не соответствующую академическим стандартам.
4. Я чувствую дискомфорт, когда использую ИИ, потому что не уверен(а), насколько точны его ответы.
5. Я беспокоюсь, что использование ИИ может привести к нечестности при выполнении учебных заданий.
6. Я испытываю тревогу, когда проверяю работы, сделанные с помощью ИИ, на предмет возможного плагиата.
7. Я опасаясь, что ИИ может генерировать материалы, нарушающие академическую честность.
8. Меня тревожит, что ИИ может непреднамеренно создавать неоригинальные тексты.
9. Я беспокоюсь о том, как ИИ влияет на честность выполнения студентами учебных заданий.
10. Я считаю, что университет предоставляет недостаточно чёткие разъяснения о том, как студентам можно использовать ИИ.
11. Я беспокоюсь из-за отсутствия единых правил преподавателей относительно применения ИИ в учебных заданиях.
12. Меня тревожит отсутствие ясной концепции, как именно студентам следует использовать ИИ при обучении.
13. Мне не хватает прозрачных правил об использовании ИИ в учебном процессе.
14. Я считаю, что официальные рекомендации по применению ИИ в обучении недостаточно разработаны.
15. Неясность требований к использованию ИИ вызывает у меня беспокойство.
16. Я опасаясь, что чрезмерное использование ИИ ухудшит мои учебные навыки.
17. Я думаю, что активное использование ИИ снижает способность студентов критически мыслить.
18. Я переживаю, что использование ИИ мешает моему развитию как самостоятельного учащегося.
19. Мне кажется, что студенты могут потерять важные навыки анализа из-за зависимости от ИИ.
20. Я опасаясь, что широкое использование ИИ в обучении снизит необходимость в самостоятельной подготовке.
21. Мне кажется, что из-за ИИ уменьшается моя мотивация самостоятельно выполнять задания.
22. Я иногда откладываю выполнение работы, думая, что ИИ поможет сделать всё быстрее.
23. Я переживаю, что ИИ может сделать меня менее уверенным(ой) в собственных знаниях.
24. Я чувствую, что чрезмерное использование ИИ снижает мою вовлеченность в учебный процесс.

Когда чаще не значит лучше: сценарии использования генеративного ИИ в высшем образовании

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-127-148

Тарасова Ксения Вадимовна – канд. пед. наук, директор Центра психометрики и измерений в образовании Института образования, ORCID: 0000-0002-3915-3165, Researcher ID: ABD-3327-2020, ktarasova@hse.ru

Талов Даниил Павлович – аспирант, младший научный сотрудник Центра психометрики и измерений в образовании Института образования, ORCID: 0000-0002-1682-0578, Researcher ID: MDS-9225-2025, dtalov@hse.ru

Иванова Алина Евгеньевна – канд. наук об образовании, заместитель директора Центра психометрики и измерений в образовании Института образования, ORCID: 0000-0003-3340-7651, Researcher ID: I-4300-2015, aeivanova@hse.ru

Национальный исследовательский университет «Высшая школа экономики», Москва, Россия
Адрес: 101000, Россия, г. Москва, Потаповский пер., 16, стр. 10

***Аннотация.** Статья посвящена анализу связи между использованием генеративного искусственного интеллекта (ИИ) в учебной деятельности и ИИ-связанными характеристиками студентов российских вузов. Исходной предпосылкой исследования является необходимость отказаться от упрощённого отождествления частоты использования ИИ с ИИ-грамотностью или компетентностью взаимодействия с генеративными системами. В статье отдельно рассматриваются четыре характеристики: функциональная опора на ИИ при принятии решений, антропоморфизация ИИ, этическая чувствительность к его использованию и навык промптинга как показатель, оцениваемый через задание с наблюдаемым продуктом. Эмпирическую базу составили данные 1647 студентов 3–4-го курсов бакалавриата из 10 российских университетов, обучающихся по STEM-дисциплинам, гуманитарным и социальным направлениям. Анализ показал, что общая частота использования генеративного ИИ положительно связана с функциональной опорой на ИИ и антропоморфизацией, но не имеет статистически значимой связи с этической чувствительностью и навыком промптинга после учёта множественных сравнений. Использование ИИ при подготовке курсовых и проектных работ оказалось связано как с функциональной опорой на ИИ, так и с антропоморфизацией. Более инструментальные сценарии, включая программирование и поиск учебных материалов, напротив, ассоциированы с более низкой антропоморфизацией. Этическая чувствительность в большей степени связана с направлением подготовки: студенты гуманитарных и социальных наук демонстрируют более высокие*

значения по сравнению со студентами STEM-направлений. Для навыка промптинга устойчивых предикторов в итоговой модели не выявлено. Полученные результаты показывают, что частота использования генеративного ИИ не может рассматриваться как достаточный индикатор ИИ-грамотности студентов. Статья обосновывает необходимость более дифференцированного подхода к обучению, оцениванию и институциональному регулированию использования ИИ в высшем образовании.

Ключевые слова: генеративный ИИ, ИИ-грамотность, высшее образование, промптинг, антропоморфизация ИИ, этика ИИ

Для цитирования: Тарасова К.В., Талов Д.П., Иванова А.Е. Когда чаще не значит лучше: сценарии использования генеративного ИИ в высшем образовании // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 127–148. – DOI: 10.31992/0869-3617-2026-35-6-127-148.

When Greater Frequency Is Not Better: Patterns of Generative AI Use in Higher Education

Original article

DOI: 10.31992/0869-3617-2026-35-6-127-148

Ksenia V. Tarasova – Cand.Sci. (PSciences), Director of the Center for Psychometrics and Measurement in Education, Institute of Education, ORCID: 0000-0002-3915-3165, Researcher ID: ABD-3327-2020, ktarasova@hse.ru

Daniil P. Talov – PhD student, Junior Research Fellow, Center for Psychometrics and Measurement in Education, Institute of Education, ORCID: 0000-0002-1682-0578, Researcher ID: MDS-9225-2025, dtalov@hse.ru

Alina E. Ivanova – Cand.Sci. (Education), Senior Research Fellow, Center for Psychometrics and Measurement in Education, Institute of Education, ORCID: 0000-0003-3340-7651, Researcher ID: I-4300-2015, aeivanova@hse.ru

National Research University Higher School of Economics, Moscow, Russian Federation
Address: 16, bldg. 10, Potapovsky lane, Moscow, 101000, Russian Federation

Abstract. The article examines the relationship between the use of generative artificial intelligence in learning and several AI-related characteristics of students at Russian universities. The study is based on the assumption that the frequency of AI use should not be treated as a direct indicator of AI literacy or competence in working with generative systems. Four characteristics are considered separately: functional reliance on AI in decision-making, anthropomorphization of AI, ethical sensitivity to its use, and prompting skill assessed through a performance-based task. The empirical analysis draws on data from 1,647 third- and fourth-year undergraduate students from 10 Russian universities, representing STEM, humanities, and social science fields. The results show that the overall frequency of generative AI use is positively associated with functional reliance on AI and anthropomorphization, but is not significantly related to ethical sensitivity or prompting skill after adjustment for multiple comparisons. The use of AI for preparing course papers and project work is associated with both functional reliance on AI and anthropomorphization. More instrumental scenarios, including programming and searching for learning materials, are associated with lower

levels of anthropomorphization. Ethical sensitivity is more strongly related to field of study: students in the humanities and social sciences demonstrate higher levels of ethical concern than students in STEM fields. No robust predictors of prompting skill were identified in the final model. The findings suggest that the frequency of generative AI use cannot serve as a sufficient indicator of students' AI literacy. The article argues for a more differentiated approach to teaching, assessment, and institutional regulation of AI use in higher education.

Keywords: generative AI, AI literacy, higher education, prompting, anthropomorphization of AI, AI-related ethics

Cite as: Tarasova, K.V., Talov, D.P., Ivanova, A.E. (2026). When Greater Frequency Is Not Better: Patterns of Generative AI Use in Higher Education. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 127-148, doi: 10.31992/0869-3617-2026-35-6-127-148 (In Russ., abstract in Eng.).

Введение

Стремительное распространение генеративного искусственного интеллекта в высшем образовании меняет не только набор цифровых инструментов, доступных студентам, но и саму организацию учебной деятельности. Генеративные модели всё чаще используются при подготовке письменных работ, поиске и переработке информации, программировании, создании презентаций, переводе, составлении библиографии и даже в ситуациях текущего контроля знаний. На этом фоне университеты сталкиваются с двойной задачей: с одной стороны, необходимо учитывать образовательный потенциал новых инструментов, связанных с повышением доступности, скорости и вариативности учебной поддержки, с другой – возрастает значимость рисков, связанных с академической добросовестностью, непрозрачностью источника ответа, перераспределением когнитивной нагрузки и ослаблением диагностической ценности традиционных учебных продуктов как свидетельств реальной компетентности студента [1–6].

Международная дискуссия вокруг генеративного ИИ в высшем образовании в последние два года заметно сместилась от общего описания возможностей технологии к обсуждению институциональных правил её использования, режимов допустимости, новых требований к оцениванию и академическому письму, а также к более сложному

вопросу – чему именно университет должен учить студента в ситуации, когда значительная часть интеллектуально насыщенных операций может быть частично делегирована модели. Исследования университетских политик показывают, что институциональные подходы к интеграции генеративного ИИ остаются крайне неоднородными – даже при наличии формальных рекомендаций роли преподавателей, студентов и администрации, а также границы допустимого использования ИИ нередко определяются неявно или противоречиво. В результате одна и та же технология может интерпретироваться либо как легитимный инструмент поддержки обучения, либо как угроза академической честности и автономии мышления [7–10].

Для российского высшего образования эта проблематика имеет особую значимость. Данные последних российских исследований показывают, что генеративный ИИ уже вошёл в повседневные образовательные практики значительной части студентов, однако наряду с инструментальным применением для ускорения учебных задач фиксируются и более глубокие изменения в представлениях студентов о допустимости ИИ-помощи, авторстве, самостоятельности и качестве результата. Российские авторы также указывают, что распространение генеративного ИИ требует пересмотра форм самостоятельной работы, процедур оценивания и способов

проверки того, что именно демонстрирует студент – собственную компетентность или способность организовать работу с внешним интеллектуальным ресурсом. Отдельной зоной напряжения становится выпускная квалификационная работа, где ИИ одновременно выступает и как средство поддержки, и как источник новых рисков для валидности итоговой оценки [11–16].

При этом одна из центральных методологических проблем современной литературы состоит не просто в недостатке данных о применении генеративного ИИ, а в смешении разных уровней анализа. В одних работах в фокусе находится поведенческая интенсивность использования, в других – установки по отношению к технологии, в третьих – риски академического мошенничества, в четвёртых – более широкий конструкт ИИ-грамотности. Методологически эти линии не являются взаимозаменяемыми: частота обращения к ИИ описывает распространённость практики, установки – её субъективную легитимацию, этическая чувствительность – нормативную рефлексию, а навык промптинга – качество конкретного действия. Когда эти измерения объединяются в один показатель «готовности» или «компетентности», возникает риск ложной каузальной интерпретации: регулярное использование начинает трактоваться как признак более развитой ИИ-грамотности, хотя оно может отражать лишь инструментальную доступность, прагматическую выгоду или снижение барьеров к делегированию учебных операций. Именно поэтому пробел современной литературы заключается не только в недостатке исследований, но и в происхождении самого пробела: он возникает из-за редукции сложного набора когнитивных, нормативных и поведенческих характеристик к наблюдаемой частоте использования [1; 17; 18].

Эта проблема особенно заметна в отношении ИИ-грамотности. Современные рамки подчёркивают, что ИИ-грамотность не сводится к техническому знакомству с инстру-

ментом или к положительному отношению к нему. В разных версиях таких рамок в неё включаются понимание принципов работы ИИ, критическая оценка его ответов, этическое и правовое измерение, социальные последствия, способность интегрировать ИИ в решение задач и рефлексивно управлять собственным взаимодействием с системой [1; 18; 19]. Для университетского контекста это означает, что вопрос о подготовке студентов к жизни и работе в среде генеративного ИИ не может решаться через простое расширение доступа к инструментам: необходима более точная операционализация тех характеристик, которые действительно следует развивать и измерять.

Особого внимания в этой связи заслуживает навык промптинга. Несмотря на его стремительную популяризацию в прикладных обсуждениях, в академической литературе он всё ещё слишком часто трактуется либо как интуитивно осваиваемая практика, либо как побочный продукт частого взаимодействия с большими языковыми моделями. Между тем исследования показывают, что формулирование эффективного запроса к генеративному ИИ требует не просто «умения написать запрос», а планирования, выделения существенных условий задачи, задания критериев желаемого результата, учёта ограничений и часто итеративной доработки взаимодействия. Более того, пользователи, не имеющие специальной подготовки, нередко строят запросы неэффективно, а качество промпта существенно связано с качеством получаемого ответа [20–22]. Следовательно, навык промптинга целесообразно рассматривать не как автоматический эффект «опыта использования ИИ», а как отдельную, поддающуюся оцениванию компетенцию.

Не менее важно и то, что использование генеративного ИИ студентами встроено в дисциплинарные контексты. Уже имеющиеся данные свидетельствуют, что представители разных направлений подготовки поразному включают ИИ в учебную деятель-

ность и по-разному осмысливают связанные с ним риски. Для *STEM*-дисциплин ИИ чаще рассматривается как инструмент повышения эффективности, решения прикладных задач и ускорения работы, тогда как в социально-гуманитарных областях сильнее представлены эпистемические, интерпретативные и этические вопросы. Сходным образом в инженерном образовании всё активнее обсуждается необходимость не внешнего, а встроенного включения этической проблематики в профессиональную подготовку, чтобы технологическая компетентность не развивалась в отрыве от нормативной рефлексии [23; 24]. Для высшего образования это означает, что одинаковые институциональные меры по интеграции ИИ могут давать неодинаковый эффект в разных дисциплинарных средах.

Таким образом, в современной литературе отчётливо обозначается несколько лакун. Во-первых, по-прежнему недостаточно исследований, которые разводили бы между собой различные ИИ-связанные характеристики студентов, а не сводили бы их к единому показателю «готовности к использованию ИИ». Во-вторых, крайне ограничены работы, в которых частота и конкретные сценарии использования генеративного ИИ сопоставляются не только с самоотчётными установками, но и с показателями, отражающими наблюдаемое качество действия. В-третьих, для российского контекста существует пока немного исследований, позволяющих одновременно рассмотреть функциональную опору на ИИ, социально-аффективное отношение к нему, этическую чувствительность и навык промптинга в одной аналитической модели. Между тем именно такое разведение необходимо, если университеты хотят строить не декларативную, а доказательную политику интеграции ИИ в обучение и оценивание.

В настоящей статье мы исходим из того, что частота использования генеративного ИИ не может автоматически рассматриваться как индикатор сформированности

ИИ-грамотности. Поэтому мы анализируем, как общая частота и конкретные учебные сценарии использования генеративного ИИ связаны с четырьмя различными ИИ-связанными характеристиками студентов: функциональной опорой на ИИ при принятии решений, антропоморфизацией ИИ, этической чувствительностью к его использованию и навыком промптинга. Такой ракурс позволяет уточнить, какие характеристики действительно сопряжены с более интенсивным использованием технологии, а какие требуют специального педагогического формирования и отдельного оценивания. Для российского высшего образования это принципиально важно, поскольку позволяет уйти от упрощённого отождествления «часто пользуется ИИ» и «компетентно работает с ИИ», а значит, точнее проектировать kurikulum, оценивание и университетские правила работы с генеративными системами.

Концептуальная рамка исследования, цель и гипотезы

Несмотря на быстрый рост числа работ, посвящённых использованию генеративного ИИ в высшем образовании, литература по-прежнему страдает от важной концептуальной неясности: в ней нередко смешиваются интенсивность использования технологии, субъективное отношение к ней, готовность делегировать ей отдельные когнитивные функции, этические установки и собственно компетентность взаимодействия с ИИ. В результате разные по своей природе феномены начинают рассматриваться как проявления одного и того же конструкта, а высокая частота использования генеративного ИИ фактически интерпретируется как индикатор сформированности ИИ-грамотности. Между тем современные рамки ИИ-грамотности, напротив, подчёркивают её многомерность и включённость в когнитивные, критические, социальные, правовые и этические измерения, не сводимые к одному показателю технологической активности [1; 18].

В настоящем исследовании мы исходим из того, что для университетского контекста особенно важно аналитически развести по меньшей мере четыре группы поведенческих характеристик студентов, связанных с использованием ИИ: функциональная опора на ИИ, то есть готовность воспринимать генеративную систему как инструмент когнитивной поддержки и принятия решений; антропоморфизация ИИ, отражающая склонность приписывать системе социально-аффективные свойства и воспринимать её не только как инструмент, но и как квазисубъекта взаимодействия; этическая чувствительность к использованию ИИ, понимаемая как выраженность нормативной озабоченности вопросами справедливости, прозрачности, безопасности, ответственности и допустимости применения ИИ в образовательной среде; навык промптинга, отражающий способность сформулировать структурированный запрос к модели с учётом условий задачи, ограничений и критериев результата.

Теоретически антропоморфизацию ИИ целесообразно позиционировать не только внутри литературы об ИИ-грамотности, но и в более широкой традиции человеко-компьютерного взаимодействия. В рамках парадигмы «Компьютеры как социальные акторы» (*Computers Are Social Actors*) задолго до массового использования ИИ было показано, что пользователи склонны применять к компьютерам социальные правила и ожидания даже тогда, когда осознают их нечеловеческую природу [25]. Сходную логику развивает трёхфакторная теория антропоморфизации, связывающая приписывание человеческих качеств нечеловеческим агентам с доступностью антропоцентрических схем, потребностью объяснять поведение объекта и мотивацией к социальному контакту [26]. Для цифровых агентов это означает, что антропоморфизация не сводится к эмоциональной симпатии к технологии: она затрагивает восприятие агентности, намерений, компетентности, доверия и ответ-

ственности. Экспериментальные работы в сфере изучения взаимодействия компьютера и человека показывают, что повышение степени человекоподобности компьютерной репрезентации усиливает социальные реакции пользователей [27], а исследования автономных систем демонстрируют, что антропоморфизация может повышать доверие и менять восприятие ответственности [28]. В более поздних работах парадигма «Компьютеры как социальные акторы» пересматривается применительно к новым формам коммуникации человек-машина, где интерактивность, умение «говорить» и персонализация делают социальные реакции на технологию ещё более вероятными [29]. Поэтому в настоящем исследовании антропоморфизация ИИ рассматривается как самостоятельный социально-когнитивный конструкт, смежный с доверием, воспринимаемой агентностью и реляционной вовлечённостью, но не тождественный ни частоте использования, ни функциональной опоре на ИИ.

Именно такое разведение позволяет избежать редукции ИИ-компетентности к частоте использования технологии и точнее описать образовательные последствия распространения генеративного ИИ. Основания для такого разведения содержатся как в современных концептуализациях ИИ-грамотности, так и в психометрических разработках, направленных на измерение различных форм зависимости, отношения и компетентности в работе с LLM [18; 30].

Функциональная опора на ИИ и антропоморфизация как разные линии включения технологии в учебную деятельность

Одним из продуктивных направлений последних исследований стало различение инструментальной и реляционной зависимости от больших языковых моделей, что отражает с одной стороны опору на ИИ в решении задач, с другой – восприятие LLM как социально значимого, «понимающего» или почти компаньоноподобного агента [30]. Такое различение особенно важно для высшего образования, поскольку позволяет не смешивать

вать делегирование части учебных операций ИИ с аффективной и антропоморфной вовлечённостью во взаимодействие с ним.

Для образовательной практики это различие принципиально. Функциональная опора на ИИ может усиливаться по мере нормализации технологии как средства ускорения поиска информации, генерации текста, структурирования ответа или поддержки решения. Однако антропоморфизация ИИ, по-видимому, зависит не столько от общей частоты контакта с системой, сколько от характера самого взаимодействия: коммуникативные, длительные и текстоцентричные сценарии могут сильнее стимулировать восприятие ИИ как «партнёра», чем короткие утилитарные обращения. Это означает, что даже при внешне сходной интенсивности использования ИИ разные студенты могут демонстрировать различные паттерны включения технологии в учебную деятельность. В более широком плане это согласуется с выводами о том, что генеративный ИИ перестраивает не только технический, но и реляционный слой образовательного опыта, включая доверие к системе, делегирование суждения и изменение границ между инструментом и квазисобеседником [1; 2].

Этическая чувствительность к ИИ как отдельная нормативная характеристика

Ещё одной проблемой является неявное предположение о том, что по мере роста опыта работы с генеративным ИИ у студентов автоматически формируется и более зрелое этическое отношение к его использованию. Однако доступные данные скорее указывают на обратное: интенсивное использование технологии само по себе не гарантирует ни понимания связанных с ней рисков, ни осознанной позиции в отношении академической добросовестности, прозрачности, авторства, ответственности и допустимых границ её применения. Международные обзоры показывают, что этические аспекты внедрения генеративного ИИ в высшем образовании требуют специальной педагогической и институциональной работы, а не возникают

как естественный побочный продукт взаимодействия с системой [2; 5; 31; 32].

Для российского высшего образования эта проблема особенно значима, поскольку отечественные публикации последних лет фиксируют одновременно и высокую скорость распространения генеративного ИИ в образовательной среде, и сохраняющуюся неопределённость его статуса в учебных и оценочных практиках. Авторы, анализирующие использование ИИ студентами российских вузов, подчёркивают различия в сценариях применения технологии, неоднозначность её влияния на самостоятельность студента и необходимость более точного нормативного и методического сопровождения её включения в образовательный процесс [13; 14]. Отсюда следует, что этическая чувствительность должна рассматриваться не как производная от интенсивности использования ИИ, а как самостоятельная нормативная характеристика, подверженная влиянию дисциплинарной культуры, образовательного контекста и институциональных норм.

Навык промптинга как наблюдаемая, а не предполагаемая компетенция

Особое место в исследовании занимает навык промптинга. В прикладных обсуждениях генеративного ИИ он часто описывается как почти естественный результат накопленного пользовательского опыта. Однако в академической литературе всё убедительнее показывается, что эффективное взаимодействие с большой языковой моделью требует отдельного комплекса действий: распознавания сути задачи, выделения релевантных условий, задания параметров желаемого результата, учёта ограничений, управления итеративностью и оценки качества отклика модели. Иными словами, промптинг представляет собой не просто технический ввод текста, а когнитивно организованное действие, связанное с планированием и эксплуатацией условий задачи [20; 21; 33].

Это имеет прямое методологическое следствие: навык промптинга не должен вы-

водиться из самоотчётов о частоте использования ИИ или из общей позитивной установки к технологии. Напротив, он требует отдельного оценивания через *performance-based* задания, в которых можно наблюдать, насколько студент способен перевести условия задачи в структурированный запрос к модели. Такой подход особенно важен в ситуации, когда сами итоговые продукты учебной деятельности всё в меньшей степени могут служить надёжным свидетельством индивидуальной компетентности, поскольку часть интеллектуальной работы может быть внешне «спрятана» внутри взаимодействия с генеративной системой [1; 2; 34].

Дисциплинарный контекст как источник различий

Наконец, имеющиеся исследования позволяют предположить, что связь студентов с генеративным ИИ опосредуется дисциплинарной средой. Работы последних лет показывают, что студенты *STEM*-дисциплин в среднем чаще рассматривают ИИ как инструмент повышения эффективности, технической поддержки и решения прикладных задач, тогда как обучающиеся в социально-гуманитарных областях чаще включают в осмысление ИИ эпистемические, интерпретативные и этические вопросы [23]. Это не означает наличия жёсткой дихотомии между направлениями подготовки, но даёт основания ожидать, что дисциплинарная принадлежность будет сильнее связана с этической чувствительностью к ИИ, чем, например, с общей функциональной опорой на него. Российские публикации также указывают, что различия в характере использования генеративного ИИ между группами студентов не исчерпываются частотой обращения к инструменту и требуют учёта образовательного профиля и типа учебных задач [14].

В совокупности эти аргументы позволяют рассматривать генеративный ИИ в высшем образовании не как единый объект «освоения», а как среду, в которой сосуществуют различные линии учебного, когнитивного, нормативного и социального взаимодей-

ствия. Поэтому эмпирическое исследование должно быть направлено не на поиск одного интегрального показателя «готовности к ИИ», а на проверку того, насколько интенсивность и сценарии использования технологии связаны с разными ИИ-связанными характеристиками студентов.

Научный вклад исследования состоит не столько в утверждении самой идеи о несводимости ИИ-грамотности к частоте использования генеративного ИИ, поскольку эта идея уже представлена в международной дискуссии, а скорее в её эмпирической проверке на данных студентов российских университетов. В отличие от работ, основанных на самоотчётных методиках об использовании ИИ или общими установками к технологии, в настоящем исследовании одновременно разведены функциональная опора на ИИ, антропоморфизация, этическая чувствительность и навык промптинга, причём последний измеряется через задание с наблюдаемым продуктом деятельности. Это позволяет показать, какие ИИ-связанные характеристики действительно сопряжены с интенсивностью и сценариями использования технологии, а какие требуют самостоятельного педагогического формирования и отдельного оценивания.

Исходя из представленной концептуальной рамки, целью исследования является анализ того, как общая частота и конкретные сценарии использования генеративного ИИ в учебной деятельности связаны с четырьмя различными ИИ-связанными характеристиками студентов российских вузов: функциональной опорой на ИИ, антропоморфизацией ИИ, этической чувствительностью к его использованию и навыком промптинга.

Для достижения этой цели были поставлены следующие исследовательские вопросы:

1) *ИВ1*. Связана ли общая частота использования генеративного ИИ с функциональной опорой на ИИ, антропоморфизацией ИИ, этической чувствительностью и навыком промптинга?

2) *ИБ2*. Различаются ли связи с указанными характеристиками в зависимости от конкретных учебных сценариев использования ИИ?

3) *ИБ3*. Как дисциплинарная принадлежность студентов связана с функциональной опорой на ИИ, антропоморфизацией ИИ, этической чувствительностью и навыком промптинга?

Концептуальная рамка исследования задаёт более строгую и методологически аккуратную логику анализа – вместо предположения о линейном переходе от частого использования ИИ к «более высокой ИИ-компетентности» мы проверяем, какие именно ИИ-связанные характеристики действительно связаны с интенсивностью и формами использования технологии, а какие требуют отдельного педагогического и оценочного внимания.

Метод

Выборка и процедура исследования

Выборка была сформирована среди университетов – участников лонгитюдного исследования «Модели образовательного поведения студентов в их связи с показателями успешности», которое проводится с 2022 года в селективных российских университетах и одобрено Комитетом по этике НИУ ВШЭ 19 сентября 2022 года. В данной работе мы опираемся на данные, собранные в 10 вузах, выборка невероятностная – университеты самостоятельно принимали решение об участии в исследовании. Затем внутри каждого укрупнённого направления подготовки (инженерное дело, технологии и технические науки; математические и естественные науки; гуманитарные науки; науки об обществе) случайным образом были отобраны учебные группы.

В каждом университете назначались представители, выполнявшие функции координаторов тестирования на месте. Они получали наборы индивидуальных учётных записей для доступа к цифровой платформе, распределяли их среди студентов, а также

информировали участников о целях и общих условиях исследования. Тестирование проводилось по заранее установленному графику в компьютерных классах университетов в стандартизированных условиях. Студенты входили на цифровую платформу с использованием уникальных логинов. После входа участникам отображался стартовый экран с основной информацией об исследовании, где подчёркивалось, что участие является добровольным, данные собираются исключительно в исследовательских целях, а участие не предполагает каких-либо вознаграждений или бонусов.

Эмпирическую базу настоящего исследования составили данные, собранные в 2025/2026 учебном году среди студентов 3–4-го курсов бакалавриата. Всего в анализ были включены данные 1647 студентов, из которых 982 обучались по *STEM*-дисциплинам, 384 – по гуманитарным направлениям, и 281 – по направлениям социальных наук.

Инструментарий

Для изучения связей между паттернами использования генеративного ИИ студентами, их когнитивными компетенциями и установками был использован комплекс инструментов, включающих самоотчётные методики и задания с оценкой продукта деятельности. Такая стратегия измерения позволяла охватить как субъективные установки по отношению к ИИ, так и наблюдаемые проявления компетентности, связанной с его использованием.

Склонность студентов к антропоморфизации ИИ-систем (реляционная зависимость) и к опоре на ИИ при принятии решений (инструментальная зависимость) оценивались с помощью локализованной и адаптированной версии шкалы *LLM-D12*, разработанной А. Янковской и соавторами [30]. Адаптированная версия включала 12 утверждений типа шкалы Ликерта и была культурно и лингвистически приведена в соответствие с контекстом исследования при сохранении исходной теоретической структуры. Инструментальная зависимость отра-

жает степень, в которой ИИ воспринимается как необходимое функциональное средство при выполнении задач и принятии решений, например, как опора на ИИ даже в тех случаях, когда человек способен справиться с задачей самостоятельно. Реляционная зависимость описывает склонность воспринимать ИИ не только как инструментальную систему, но и как квазисоциального участника взаимодействия, что может проявляться в антропоморфизации ИИ, приписывании ему признаков понимания, поддержки или эмоциональной отзывчивости.

Такое различие позволяет эмпирически разграничить прагматическую опору на ИИ как инструмент и восприятие его как личности. Это разграничение особенно важно для понимания доверия к ИИ, делегирования ему суждений и восприятия его как действующего агента в образовательной среде.

Также была разработана шкала (в формате шкалы Ликерта), предназначенная для оценки нормативных и ценностных ориентаций студентов в отношении ИИ – чувствительность к вопросам справедливости, прозрачности, безопасности и ответственности при разработке и применении ИИ-систем. Формулировки утверждений отражали согласие с общими нормативными принципами, например с необходимостью учитывать этические последствия использования или поддерживать регулирование ИИ, при этом утверждения не проверяли знания конкретных норм или формальных правил. Такой подход позволял рассматривать этическую чувствительность как относительно устойчивую установку, а не как показатель технической осведомлённости и знания нормативных требований.

Навык промптинга был операционализирован как показатель компетентности в использовании инструментов ИИ, отражающий способность формулировать ясные, структурированные и целенаправленные запросы к большим языковым моделям. Концептуально промптинг рассматривался не как простой ввод текста, а как когнитивно

нагруженная деятельность, требующая планирования, абстрагирования условий задачи и явного задания критериев успешного результата [34].

Участникам предлагалось составить промпт объёмом до 300 слов для большой языковой модели с целью генерации рекомендаций по развитию конкретного бизнеса. В инструкции к заданию содержалось краткое пояснение того, что понимается под промптом, а также примеры, прямое копирование текста задания запрещалось. Задание включало 10 ключевых условий, встроенных в описание ситуации. Каждое условие оценивалось дихотомически (0 – отсутствует, 1 – присутствует) в зависимости от того, было ли оно явно отражено в тексте промпта участника. Оценивание проводилось по ранее апробированному протоколу: автоматизированное кодирование по аналитической рубрике было проверено двумя экспертами в области педагогических измерений на 10-процентной подвыборке ответов, согласованность между экспертными оценками и оценками модели была высокой (средняя точность = 0,93; Коэна = 0,80) [34].

Такой способ оценивания обеспечивал объективность и снижал влияние избыточности или стиля изложения, позволяя сосредоточиться на структурных и функциональных характеристиках навыка.

Помимо основных шкал и задания для оценки навыков промптинга, использовалась анкета, предназначенная для фиксации поведенческих, когнитивных и институциональных паттернов использования генеративного ИИ. В этот блок вошли показатели частоты использования ИИ, типов учебных задач, решаемых с его помощью, целей и сценариев использования, восприятия институциональной политики, взаимодействия с преподавателями по поводу ИИ, опыта формального обучения, связанного с ИИ, а также уровня самостоятельного освоения этих технологий. Эти переменные использовались в качестве предикторов в последующем регрессионном анализе.

Стратегия анализа

Аналитический подход был специально выстроен таким образом, чтобы разделить эффекты общей интенсивности использования генеративного ИИ и использования ИИ в конкретных учебных задачах, с учётом различий между академическими направлениями и обеспечением надёжности статистических выводов.

Четыре зависимые переменные (опора на ИИ при принятии решений, антропоморфизация ИИ, этическая чувствительность и навык промптинга) оценивались в виде факторных баллов, полученных на основе факторного анализа с порядковыми показателями с использованием метода оценки *WLSMV (Weighted Least Squares Mean and Variance adjusted)*. Качество модели было высоким: $CFI = 0,986$, $TLI = 0,985$, $RMSEA = 0,050$, $SRMR = 0,057$. Полученные факторные баллы были стандартизованы (z -преобразование) перед проведением регрессионного анализа.

В модели были включены три группы предикторов.

1. Частота использования ИИ в учебных целях

Общая частота использования ИИ определялась с помощью одного вопроса, в котором респондентов спрашивали, как часто они используют ИИ в учебных целях. Исходная шкала ответов варьировала от 1 («совсем не использую») до 5 («использую каждый день»). Ответ «Затрудняюсь ответить» рассматривался как пропущенное значение. Далее переменная была стандартизована с использованием z -преобразования и включена в модели как непрерывный предиктор.

2. Task-specific использование ИИ

Для учёта разнообразия практик использования ИИ был сформирован набор бинарных показателей (0 – не выбрано, 1 – выбрано) на основе вопроса с множественным выбором о том, применяют ли студенты ИИ для различных типов учебных задач. В качестве отдельных предикторов рассматривались следующие категории:

- генерация текстов для письменных заданий;
- генерация текстов для курсовых и проектных работ;
- поиск и обработка учебных материалов с помощью ИИ;
- подготовка устных выступлений с использованием ИИ;
- перевод текстов с помощью ИИ;
- использование ИИ в практических и лабораторных работах;
- суммирование текстов с помощью ИИ;
- программирование с использованием ИИ;
- помощь ИИ во время тестов и экзаменов;
- работа с источниками и библиографией с помощью ИИ.

Такая спецификация позволила рассматривать использование ИИ как совокупность различных учебных практик, а не как единое обобщённое поведение.

3. Направление подготовки

Направление подготовки включалось в модель как категориальный предиктор с тремя уровнями: естественнонаучное направление (далее *STEM*, референтная категория), гуманитарные направления и социальные науки.

Потенциальный модерационный эффект академического направления в связях между предикторами использования ИИ и зависимыми переменными проверялся с помощью включения взаимодействий между направлением подготовки и ключевыми предикторами.

Однако добавление этих взаимодействий не привело к статистически значимому улучшению качества моделей. В связи с отсутствием дополнительной объяснительной и предсказательной ценности взаимодействия были исключены из окончательных моделей, и в дальнейшем анализе учитывались только основные эффекты (Табл. 1). Такое решение было принято исходя из соображений парсимонии и интерпретируемости моделей.

Все модели оценивались с использованием линейной регрессии. Проверки допу-

щений выявили нарушение гомоскедастичности, в связи с чем были использованы робастные стандартные ошибки, устойчивые к гетероскедастичности (*НСЗ*). По сравнению со стандартными оценками метода наименьших квадратов, *НСЗ* обеспечивает более надёжные проверки значимости коэффициентов и значения p при наличии неоднородности дисперсий и рекомендуется для моделей со множеством предикторов.

С учётом относительно большого числа предикторов в каждой модели применялась корректировка уровня значимости по процедуре контроля ложного обнаружения (эффекта) при множественной проверке гипотез Бенджамини – Хохберга в пределах каждой зависимой переменной. Для оценки статистической значимости использовались скорректированные значения p (qBH), что позволило контролировать долю ложноположительных результатов при сохранении достаточной чувствительности к существенным эффектам.

Результаты

Во всех построенных моделях доля объяснённой дисперсии оказалась умеренной. Модели объясняли 8,8% дисперсии показателя опоры на ИИ при принятии решений, 5,3% – антропоморфизации ИИ, 4,0% – этики, и 2,5% – навыка промптинга ($n = 1647$).

Опора на ИИ при принятии решений

Наиболее сильным и стабильным предиктором оказалась общая частота использования ИИ в учебной деятельности: она положительно связана с показателем опоры на ИИ при принятии решений ($\beta = 0,218$, $qBH < 0,001$). Кроме того, использование ИИ для генерации текстов при выполнении курсовых и проектных работ также было положительно связано с этим показателем ($\beta = 0,259$, $qBH = 0,0013$). Более слабый, но статистически подтверждённый эффект наблюдался для использования ИИ при подготовке устных выступлений ($\beta = 0,129$, $qBH = 0,0458$), при этом его значимость находилась на границе при более строгой глобальной корректировке.

После применения поправки Бенджамини – Хохберга в пределах данной зависимой переменной академическое направление не продемонстрировало статистически надёжной связи с принятием решений с опорой на ИИ. Хотя коэффициент для студентов гуманитарных направлений по сравнению с естественнонаучным направлением был отрицательным ($b = -0,129$), этот эффект не сохранил значимости после учёта множественных сравнений.

Социальное одушевление (антропоморфизация) ИИ

В модели, предсказывающей антропоморфизацию ИИ, несколько предикторов оказались статистически значимыми. Частота использования ИИ положительно связана с антропоморфизацией ($\beta = 0,160$, $qBH < 0,001$). Аналогичная положительная связь наблюдалась для использования ИИ при генерации текстов в рамках курсовых и проектных работ ($\beta = 0,245$, $qBH = 0,0052$).

В то же время более инструментальные формы использования ИИ были отрицательно связаны с антропоморфизацией. В частности, использование ИИ для задач, связанных с программированием, ассоциировано с более низкими показателями антропоморфизации ($\beta = -0,147$, $qBH = 0,0259$), как и использование ИИ для поиска и обработки учебных материалов ($\beta = -0,137$, $qBH = 0,0259$).

Академическое направление также показало статистически подтверждённую связь: студенты гуманитарных направлений демонстрировали более низкий уровень антропоморфизации ИИ по сравнению со студентами естественнонаучных направлений ($b = -0,159$, $p = 0,0356$). При учёте поправки на множественное сравнение этот эффект становился пограничным ($qBH = 0,0594$). Другие показатели использования ИИ в различных типах задач не достигли статистической значимости после поправки.

Этика, связанная с использованием ИИ

Для показателя этической чувствительности наиболее выраженные эффекты были

Таблица 1

Предикторы, связанные с использованием ИИ в различных типах учебных задач, и их вклад в показатели, связанные с ИИ (стандартизированные коэффициенты)

Table 1

Predictors related to AI use across different types of academic tasks and their contribution to AI-related indicators (standardized coefficients)

Предиктор	Принятие решений (β)	Антропоморфизация (β)	Этика (β)	Навык промптинга (β)
Общая частота использования ИИ	0,22***	0,16***	0,05	-0,02
Генерация текстов для письменных заданий	0,05	0,03	-0,01	-0,03
Генерация текстов для курсовых и выпускных работ	0,26**	0,25**	-0,13	-0,17
Поиск и обработка учебных материалов с помощью ИИ	-0,06	-0,14*	0,11	0,07
Подготовка устных выступлений с использованием ИИ	0,13*	0,09	0,03	0,03
Перевод текстов с помощью ИИ	0,03	0,07	-0,06	-0,03
Использование ИИ в практических и лабораторных работах	0,07	0,02	-0,02	0,11
Суммирование текстов с помощью ИИ	-0,03	-0,02	0,03	0,14
Программирование с использованием ИИ	-0,09	-0,15*	0,07	0,12
ИИ как помощник во время тестов и экзаменов	0,06	0,08	-0,03	-0,03
Работа с источниками и библиографией с помощью ИИ	-0,06	-0,08	0,01	0,13
Гуманитарные науки (по сравнению со STEM)	-0,13	-0,16*	0,37***	0,11
Социальные науки (по сравнению со STEM)	-0,05	-0,05	0,26***	0,15
RI	0,088	0,053	0,040	0,025

Примечание: стандартизированные коэффициенты регрессии (β) из линейных регрессионных моделей. Использованы робастные стандартные ошибки, устойчивые к гетероскедастичности (HC3). Значения p скорректированы с использованием процедуры контроля ложного открытия Бенджамини – Хохберга в пределах каждой зависимой переменной. Категория STEM используется как базовая для академической дисциплины. * $p < 0,05$, ** $p < 0,01$, *** $p < 0,001$.

Note: Standardized regression coefficients (β) from linear regression models are reported. Robust standard errors resistant to heteroskedasticity (HC3) were used. p values were adjusted using the Benjamini - Hochberg false discovery rate procedure within each dependent variable. The STEM category was used as the reference category for academic discipline. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$.

связаны с направлением подготовки. По сравнению со студентами STEM-дисциплин, студенты гуманитарных направлений демонстрировали существенно более высокий уровень этической чувствительности ($\beta = 0,369$, $qBH < 0,001$). Студенты социальных наук также показывали более высокие значения

по сравнению с естественнонаучной группой ($\beta = 0,260$, $qBH = 0,0002$).

Другие предикторы, включая использование ИИ в конкретных типах задач, в ряде случаев имели ожидаемое направление эффектов (например, использование ИИ для поиска и обработки информации), однако не

достигали статистической значимости после поправки Бенджамини – Хохберга.

Навык промптинга

Для навыка промптинга ряд предикторов демонстрировал значимость на уровне нескорректированных p -значений, включая использование ИИ для аннотирования текстов и обучение по направлениям социальных наук. Однако после учёта множественных сравнений ни один из эффектов не сохранил статистической значимости. Минимальные скорректированные значения p составляли около 0,11, что указывает на отсутствие устойчивых предикторов уровня владения этим навыком в рамках текущей модели.

Как видно из таблицы 1, общая частота использования ИИ положительно связана с опорой на ИИ при принятии решений ($\beta = 0,22, p < ,001$) и с очеловечиванием ИИ ($\beta = 0,16, p < ,001$), но не демонстрирует статистически значимой связи ни с этической озабоченностью ($\beta = 0,05$), ни с навыком промптинга ($\beta = -0,02$). Среди предикторов, связанных с деятельностью, наиболее выраженные положительные связи наблюдаются для генерации текста для курсовых и проектных работ: с опорой на ИИ при принятии решений ($\beta = 0,26, p < ,01$) и с антропоморфизацией ИИ ($\beta = 0,25, p < ,01$). Напротив, использование ИИ в задачах, связанных с программированием ($\beta = -0,15, p < ,05$) и использование ИИ для поиска и обработки учебных материалов ($\beta = -0,14, p < ,05$) отрицательно связаны с очеловечиванием ИИ. Для этической озабоченности значимыми оказываются только различия по направлениям подготовки: гуманитарные направления ($\beta = 0,37, p < ,001$) и социальные науки ($\beta = 0,26, p < ,001$) по сравнению со *STEM*. Для навыка промптинга статистически значимых предикторов в итоговой модели не выявлено.

Дискуссия

Относительно ответа на первый поставленный нами исследовательский вопрос полученные результаты показывают, что

более частое использование ИИ в учебной деятельности связано как с большей опорой на ИИ при принятии решений, так и с более выраженной антропоморфизацией ИИ. Этот результат в целом согласуется с современными исследованиями, показывающими, что студенты всё чаще включают генеративные ИИ-инструменты в повседневную учебную работу и воспринимают их как средство ускорения, структурирования и облегчения выполнения учебных задач [3; 6]. Вместе с тем наши данные позволяют уточнить эту картину. Частое использование ИИ связано не только с инструментальной нормализацией технологии, но и с большей готовностью опираться на неё как на источник суждения, подсказки или решения. Иными словами, речь идёт не просто о расширении набора цифровых средств, а о постепенном изменении того, как студент распределяет когнитивную работу между собой и внешней интеллектуальной системой.

В то же время доля объяснённой дисперсии во всех моделях остаётся ограниченной, и этот результат важен для интерпретации полученных связей. Его не следует трактовать как простое свидетельство «слабости» моделей, однако он показывает, что частота и сценарии использования ИИ объясняют лишь часть вариативности ИИ-связанных характеристик студентов. Такая величина объяснённой дисперсии выглядит сопоставимой с результатами исследований, где анализируются не только намерение использовать ИИ, но и более сложные характеристики пользовательского опыта, уверенности или ответственного применения генеративного ИИ [35; 36]. Это особенно важно для основной логики исследования – использование генеративного ИИ не тождественно сформированности ИИ-грамотности или компетентности взаимодействия с ним. Установки, этическая чувствительность и навыки работы с ИИ, по-видимому, формируются под влиянием более широкого набора факторов: особенностей образовательной среды, характера учебных заданий,

дисциплинарной культуры, эпистемических установок студентов, опыта обратной связи и уровня метакогнитивной регуляции, самоэффективности студентов, их познавательного интереса [35; 36].

Отдельного внимания заслуживает результат, связанный с навыком промптинга. Если функциональная опора на ИИ и антропоморфизация в большей степени связаны с частотой и сценариями использования технологии, то навык промптинга не должен интерпретироваться как прямое следствие самого факта обращения к ИИ. Его измерение через задание с наблюдаемым продуктом позволяет увидеть принципиально иной аспект взаимодействия с генеративной моделью – не то, насколько часто студент использует ИИ и как он к нему относится, а насколько он способен перевести условия задачи в структурированный запрос, задать ограничения, критерии результата и необходимые параметры ответа. Именно этот результат является методологически значимым, он показывает, что пользовательская активность и качество действия с ИИ находятся в разных плоскостях измерения. Для образовательного оценивания это означает, что самоотчёты о частоте использования ИИ не могут заменять оценку реальных навыков работы с ним.

Результаты также позволяют выделить различные паттерны использования ИИ, отвечая на наш второй исследовательский вопрос. В частности, положительные и статистически значимые связи с антропоморфизацией ИИ обнаружены для более длительных и коммуникативно насыщенных учебных задач – подготовки курсовых и выпускных работ, а также помощи в подготовке устных выступлений. Напротив, более инструментальные формы использования ИИ, например помощь при программировании или поиск информации, связаны с более низким уровнем антропоморфизации. Это даёт возможность предположить, что реляционный компонент взаимодействия с ИИ формируется не столько под влиянием самой частоты использования, сколько под влиянием

характера задачи. Там, где студент вступает с моделью в развёрнутое текстовое взаимодействие, обращается к ней за поддержкой аргументации, структурированием мысли или подготовкой публичного высказывания, ИИ с большей вероятностью начинает восприниматься как квазисоциальный участник учебной деятельности. Напротив, краткие утилитарные обращения, связанные с поиском информации или технической помощью, в меньшей степени поддерживают такую антропоморфную интерпретацию. Такая дифференциация согласуется с результатами обзоров, показывающих значительные различия в реакции студентов в зависимости от типа задач и учебного контекста [4].

Наконец, пытаясь ответить на наш третий исследовательский вопрос, мы показали, что роль направления подготовки оказывается более выраженной для этической оценки использования ИИ, чем для функциональных аспектов его применения. Ранее было показано, что студенты осознают потенциальные риски использования ИИ, включая плагиат, академическую недобросовестность и возможное снижение уровня критического мышления [3; 4]. Наши результаты уточняют этот вывод – этическая чувствительность распределена неравномерно и, по-видимому, связана с дисциплинарным контекстом, в котором студенты осваивают способы работы с текстом, знанием, авторством и доказательностью. Студенты гуманитарных и социальных направлений демонстрируют более высокий уровень этической чувствительности по сравнению со студентами *STEM*-направлений. Это не следует трактовать как дефицит одной группы студентов или преимущество другой. Скорее, результат указывает на то, что этика не формируется автоматически вместе с техническим или инструментальным опытом использования ИИ. Поэтому интеграция ИИ в образовательные программы должна включать не только обучение применению инструментов, но и явное обсуждение границ допустимой помощи, авторства, ответ-

ственности, прозрачности и академической добросовестности. Особенно важно это для тех образовательных контекстов, где ИИ преимущественно осмысливается как средство повышения эффективности решения прикладных задач.

В совокупности полученные результаты подтверждают центральный тезис исследования — частота использования генеративного ИИ не может рассматриваться как универсальный индикатор компетентности или ИИ-грамотности студента. Этот вывод согласуется с новыми эмпирическими данными, показывающими, что даже при высокой распространённости генеративного ИИ среди студентов самооценка навыков промптинга, способность критически оценивать ответы модели и этическая рефлексия распределены неравномерно и не сводятся к самому факту использования технологии [36]. Более интенсивное использование технологии действительно связано с большей функциональной опорой на ИИ и, в определённых сценариях, с антропоморфизацией системы. При этом сценарий использования принципиально важен: качественные исследования студенческого опыта показывают, что генеративные чат-боты включаются в разные учебные практики от поиска идей и языковой доработки до подготовки академических текстов, но именно в этих практиках особенно остро проявляются вопросы надёжности результата, авторства, плагиата и академических норм [37]. Этическая чувствительность и навык промптинга имеют иной статус и требуют самостоятельного педагогического и измерительного внимания. Особенно это касается промптинга, современные исследования рассматривают его не как естественный побочный продукт частого обращения к ИИ, а как отдельное когнитивно организованное действие, связанное со структурированием условий задачи, заданием ограничений, критериев результата и управлением взаимодействием с моделью [22; 33]. Для университетов это означает, что политика интеграции ИИ не должна ограни-

чиваться разрешением или запретом отдельных инструментов. Более продуктивным является переход к дифференцированной модели — отдельно оценивать, как студенты используют ИИ, как они понимают его ограничения, насколько осознают этические риски и способны ли качественно организовать взаимодействие с моделью. Именно такая логика позволяет связать обсуждение генеративного ИИ не только с академической честностью, но и с задачами обновления образовательного оценивания в высшей школе.

Заключение

В исследовании было показано, что частота использования генеративного ИИ в учебной деятельности связана прежде всего с функциональной опорой на ИИ и, в меньшей степени, с антропоморфизацией системы. При этом этическая чувствительность и навык промптинга не сводятся к интенсивности использования технологии. Особенно важно, что навык промптинга после учёта множественных проверок не продемонстрировал устойчивой связи ни с общей частотой обращения к ИИ, ни с отдельными сценариями его применения. Это подтверждает ключевой тезис статьи о том, что пользовательская активность не может рассматриваться как надёжный индикатор компетентного взаимодействия с генеративными системами.

Полученные результаты также показывают, что значение имеет не только частота, но и характер учебных задач, в которых используется ИИ. Коммуникативно насыщенные и текстоцентричные сценарии, прежде всего подготовка курсовых и проектных работ, связаны с более выраженной антропоморфизацией, тогда как инструментальные формы использования, включая программирование и поиск информации, напротив, ассоциированы с её снижением. Этическая чувствительность в большей степени связана с направлением подготовки. Студенты гуманитарных и социальных наук демонстрируют более высокий уровень этической чувствительности по сравнению со студентами

STEM-направлений. Эти различия указывают на то, что ИИ-связанные характеристики студентов формируются не только через опыт использования технологии, но и в контексте дисциплинарных норм, типов учебных заданий и способов работы со знанием.

Ограничения исследования связаны с тем, что анализ основан на одном срезе данных, поэтому результаты следует интерпретировать как ассоциации, а не как свидетельство причинного влияния использования ИИ на установки и навыки студентов. В дальнейшем важно анализировать динамику ИИ-связанных характеристик студентов и учитывать роль университетских правил, преподавательских требований и обратной связи. В практическом плане результаты показывают, что университетская политика в области ИИ должна выходить за рамки разрешения или запрета инструментов и опираться на более дифференцированную модель, в частности, отдельно формировать и оценивать навыки промптинга, понимание ограничений ИИ, этическую рефлексию и способность студента сохранять собственную интеллектуальную позицию.

Конфликт интересов

Авторы заявляют об отсутствии конфликта интересов. Исследователи не участвовали в преподавании дисциплины, а также в проведении текущей и итоговой аттестации студентов, чьи данные использовались в анализе.

Исследование проводилось в рамках лонгитюдного проекта «Модели образовательного поведения студентов в их связи с показателями успешности», одобренного Комитетом по этике НИУ ВШЭ 19 сентября 2022 г. Участие студентов было добровольным, данные собирались и анализировались в обезличенном виде и использовались исключительно в исследовательских целях.

Литература

1. Miao F., Holmes W. Guidance for generative AI in education and research. – Paris: UNESCO, 2023. – 44 p. – DOI: 10.54675/EWZM9535.
2. OECD. Trends shaping education 2025. – Paris: OECD Publishing, 2025. – 101 p. – DOI: 10.1787/ee6587fd-en
3. Sousa A.E., Cardoso P. Use of generative AI by higher education students // *Electronics*. – 2025. – Vol. 14, no. 7. – Article no. 1258. – DOI: 10.3390/electronics14071258.
4. Wu F., Dang Y., Li M. A systematic review of responses, attitudes, and utilization behaviors on generative AI for teaching and learning in higher education // *Behavioral Sciences*. – 2025. – Vol. 15, no. 4. – Article no. 467. – DOI: 10.3390/bs15040467.
5. Bittle K., El-Gayar O. Generative AI and academic integrity in higher education: A systematic review and research agenda // *Information*. – 2025. – Vol. 16, no. 4. – Article no. 296. – DOI: 10.3390/info16040296.
6. Çerkinî B., Bajraktari K., Çibuçiu B., Ramadani F., Zejnullahu F., Hajdini L. Artificial intelligence in higher education: Student perspectives on practices, challenges, and policies in a transitional context // *Frontiers in Education*. – 2025. – Vol. 10. – Article no. 1700056. – DOI: 10.3389/feduc.2025.1700056.
7. Jin Y., Yan L., Echeverria V., Gašević D., Martinez-Maldonado R. Generative AI in higher education: A global perspective of institutional adoption policies and guidelines // *Computers and Education: Artificial Intelligence*. – 2025. – Vol. 8. – Article no. 100348. – DOI: 10.1016/j.caeai.2024.100348.
8. Aristombayeva M., Satybaldiyeva R., Maung B.M., Lee D. Guiding the uncharted: The emerging (and missing) policies on generative AI in higher education // *Frontiers in Education*. – 2025. – Vol. 10. – Article no. 1644081. – DOI: 10.3389/feduc.2025.1644081.
9. Tsao J. Trajectories of AI policy in higher education: Interpretations, discourses, and enactments of students and teachers // *Computers and Education: Artificial Intelligence*. – 2025. – Vol. 9. – Article no. 100496. – DOI: 10.1016/j.caeai.2025.100496.
10. Miao F., Shiohira K., Lao N. AI competency framework for students. – Paris: UNESCO, 2024. – 80 p. – DOI: 10.54675/JKJB9835.
11. Алешковский И.А., Гаспаршвили А.Т., Нарбут Н.П., Крухмалева О.В., Савина Н.Е. Российские студенты о возможностях и ограничениях искусственного интеллекта в обучении // *Вестник РУДН. Серия: Социоло-*

- гия. – 2024. – Т. 24, № 2. – С. 335–353. – DOI: 10.22363/2313-2272-2024-24-2-335-353.
12. Буякова К.И., Дмитриев Я.А., Иванова А.С., Фещенко А.В., Яковлева К.И. Отношение студентов и преподавателей к использованию инструментов с искусственным интеллектом в вузе // Образование и наука. – 2024. – Т. 26, № 7. – С. 160–193. – DOI: 10.17853/1994-5639-2024-7-160-193.
 13. Кошкина Е.А., Бордовская Н.В., Гнедых Д.С., Хромова М.А., Демьянчук Р.В., Исакова М.П., Балышев П.А. Генеративный искусственный интеллект в высшем образовании: обзор теоретических подходов и практик применения // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 36–57. – DOI: 10.31992/0869-3617-2025-34-6-36-57.
 14. Кузьминов Я.И., Кручинская Е.В., Груздев И.А., Наумов А.А. Отстающие и опережающие: как студенты используют генеративный искусственный интеллект в образовательных целях // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 9–35. – DOI: 10.31992/0869-3617-2025-34-6-9-35.
 15. Тихонова Н.В., Поморцева Н.П. Выпускная квалификационная работа в вузе в условиях распространения искусственного интеллекта: взгляд студентов // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 112–135. – DOI: 10.31992/0869-3617-2025-34-6-112-135.
 16. Земцов Д.И., Груздев И.А. «Цифровой кентавр»: совместное обучение человека и ИИ в университете // Высшее образование в России. – 2025. – Т. 34, № 10. – С. 47–62. – DOI: 10.31992/0869-3617-2025-34-10-47-62.
 17. Oncioiu I., Bularca A. R. Artificial intelligence governance in higher education: The role of knowledge-based strategies in fostering legal awareness and ethical artificial intelligence literacy // Societies. – 2025. – Vol. 15, no. 6. – Article no. 144. – DOI: 10.3390/soc15060144.
 18. Hackl V., Müller A.E., Sailer M. The AI literacy heptagon: A structured approach to AI literacy in higher education // Computers and Education: Artificial Intelligence. – 2026. – Vol. 10. – Article no. 100540. – DOI: 10.1016/j.caeai.2026.100540.
 19. Zhang C., Magerko B. Generative AI literacy: A comprehensive framework for literacy and responsible use. – arXiv:2504.19038. – 2025. – DOI: 10.48550/arXiv.2504.19038.
 20. Zamfirescu-Pereira J.D., Wong R.Y., Hartmann B., Yang Q. Why Johnny can't prompt: How non-AI experts try (and fail) to design LLM prompts // Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems. – New York: Association for Computing Machinery, 2023. – Article no. 437. – P. 1–21. – DOI: 10.1145/3544548.3581388.
 21. Federiakin D., Molerov D., Zlatkin-Troitschanskaia O., Maur A. Prompt engineering as a new 21st century skill // Frontiers in Education. – 2024. – Vol. 9. – Article no. 1366434. – DOI: 10.3389/educ.2024.1366434.
 22. Carrasco-Sáez J.L., Contreras-Saavedra C., San-Martín-Quiroga S., Contreras-Saavedra C.E., Viveros-Muñoz R. Analyzing higher education students' prompting techniques and their impact on ChatGPT's performance: An exploratory study in Spanish // Applied Sciences. – 2025. – Vol. 15, no. 14. – Article no. 7651. – DOI: 10.3390/app15147651.
 23. Qu Y., Tan M.X.Y., Wang J. Disciplinary differences in undergraduate students' engagement with generative artificial intelligence // Smart Learning Environments. – 2024. – Vol. 11. – Article no. 51. – DOI: 10.1186/s40561-024-00341-6.
 24. Ferdman A., Ratti E. What do we teach to engineering students: Embedded ethics, morality, and politics // Science and Engineering Ethics. – 2024. – Vol. 30. – Article no. 7. – DOI: 10.1007/s11948-024-00469-1.
 25. Nass C., Moon Y. Machines and mindlessness: Social responses to computers // Journal of Social Issues. – 2000. – Vol. 56, no. 1. – P. 81–103. – DOI: 10.1111/0022-4537.00153.
 26. Epley N., Waytz A., Cacioppo J.T. On seeing human: A three-factor theory of anthropomorphism // Psychological Review. – 2007. – Vol. 114, no. 4. – P. 864–886. – DOI: 10.1037/0033-295X.114.4.864.
 27. Gong L. How social is social responses to computers? The function of the degree of anthropomorphism in computer representations // Computers in Human Behavior. – 2008. – Vol. 24, no. 4. – P. 1494–1509. – DOI: 10.1016/j.chb.2007.05.007.
 28. Waytz A., Heafner J., Epley N. The mind in the machine: Anthropomorphism increases trust in an autonomous vehicle // Journal of Experimental Social Psychology. – 2014. – Vol. 52. – P. 113–117. – DOI: 10.1016/j.jesp.2014.01.005.
 29. Gambino A., Fox J., Ratan R.A. Building a stronger CASA: Extending the computers are social actors paradigm // Human-Machine Communication. – 2020. – Vol. 1. – P. 71–86. – DOI: 10.30658/hmc.1.5.

30. Yankouskaya A., Babiker A.B., Rizvi S.W.F., Alshakhsi S., Liebherr M., Ali R. LLM-D12: A dual-dimensional scale of instrumental and relational dependencies on large language models // *ACM Transactions on the Web*. – 2025. – Advance online publication. – DOI: 10.1145/3765895.
 31. Pikhart M., Al-Obaydi L.H. Reporting the potential risk of using AI in higher education: Subjective perspectives of educators // *Computers in Human Behavior Reports*. – 2025. – Vol. 18. – Article no. 100693. – DOI: 10.1016/j.chbr.2025.100693.
 32. Huang D., Hash N., Cummings J.J., Prena K. Academic cheating with generative AI: Exploring a moral extension of the theory of planned behavior // *Computers and Education: Artificial Intelligence*. – 2025. – Vol. 8. – Article no. 100424. – DOI: 10.1016/j.caeai.2025.100424.
 33. Pangestu F.N.N., Hamidah I., Widaningsih L., Abdullah A.G., Saputra N.A. A scoping literature review of prompt engineering for bridging students' AI literacy in higher education // *Computers and Education: Artificial Intelligence*. – 2026. – Vol. 10. – Article no. 100581. – DOI: 10.1016/j.caeai.2026.100581.
 34. Иванова А.Е., Тарасова К.В., Талов Д.П. Между интересом и умением: как студенты воспринимают и применяют ИИ // *Высшее образование в России*. – 2025. – Т. 34, № 8–9. – С. 9–32. – DOI: 10.31992/0869-3617-2025-34-8-9-32.
 35. Consoli T., Petko D. Which educational approaches predict students' generative AI confidence and responsibility? // *Computers and Education: Artificial Intelligence*. – 2025. – Vol. 9. – Article no. 100431. – DOI: 10.1016/j.caeai.2025.100431.
 36. Schutte N.S., Li H. The role of self-efficacy and curiosity in student use of artificial intelligence (AI) // *International Journal of Educational Technology in Higher Education*. – 2025. – Vol. 22. – Article no. 73. – DOI: 10.1186/s41239-025-00574-6.
 37. Jongkind R., Bikker Y., Meinema J., Broens T. Studying with GenAI: Cross-sectional study on usage patterns, needs, competencies, and ethical perspectives of medical informatics students // *Frontiers in Education*. – 2025. – Vol. 10. – Article no. 1658415. – DOI: 10.3389/educ.2025.1658415.
 38. Hadinejad N., Sperling K., McGrath C. Generative AI chatbots in higher education: Student experiences and perceived ethical challenges // *Computers and Education Open*. – 2025. – Vol. 9. – Article no. 100311. – DOI: 10.1016/j.caeo.2025.100311.
- Благодарности.** Исследование осуществлено в рамках Программы фундаментальных исследований НИУ ВШЭ (HSE-BR-2025-009).
- Статья поступила в редакцию 29.04.2026
Принята к публикации 04.06.2026

References

1. Miao, F., Holmes, W. (2023). *Guidance for Generative AI in Education and Research*. Paris : UNESCO, 44 p., doi: 10.54675/EWZM9535.
2. OECD. (2025). *Trends Shaping Education 2025*. Paris : OECD Publishing, 101 p., doi: 10.1787/ee6587fd-en.
3. Sousa, A.E., Cardoso, P. (2025). Use of Generative AI by Higher Education Students. *Electronics*. Vol. 14, no. 7, article no. 1258, doi: 10.3390/electronics14071258.
4. Wu, F., Dang, Y., Li, M. (2025). A Systematic Review of Responses, Attitudes, and Utilization Behaviors on Generative AI for Teaching and Learning in Higher Education. *Behavioral Sciences*. Vol. 15, no. 4, article no. 467, doi: 10.3390/bs15040467.
5. Bittle, K., El-Gayar, O. (2025). Generative AI and Academic Integrity in Higher Education: A Systematic Review and Research Agenda. *Information*. Vol. 16, no. 4, article no. 296, doi: 10.3390/info16040296.
6. Çerkini, B., Bajraktari, K., Çibuçiu, B., Ramadani, F., Zejnullahu, F., Hajdini, L. (2025). Artificial Intelligence in Higher Education: Student Perspectives on Practices, Challenges, and Policies in a Transitional Context. *Frontiers in Education*. Vol. 10, article no. 1700056, doi: 10.3389/educ.2025.1700056.

7. Jin, Y., Yan, L., Echeverria, V., Gašević, D., Martinez-Maldonado, R. (2025). Generative AI in Higher Education: A Global Perspective of Institutional Adoption Policies and Guidelines. *Computers and Education: Artificial Intelligence*. Vol. 8, article no. 100348, doi: 10.1016/j.caeai.2024.100348.
8. Aristombayeva, M., Satybaldiyeva, R., Maung, B.M., Lee, D. (2025). Guiding the Uncharted: The Emerging (and Missing) Policies on Generative AI in Higher Education. *Frontiers in Education*. Vol. 10, article no. 1644081, doi: 10.3389/educ.2025.1644081.
9. Tsao, J. (2025). Trajectories of AI Policy in Higher Education: Interpretations, Discourses, and Enactments of Students and Teachers. *Computers and Education: Artificial Intelligence*. Vol. 9, article no. 100496, doi: 10.1016/j.caeai.2025.100496.
10. Miao, F., Shiohira, K., Lao, N. (2024). *AI Competency Framework for Students*. Paris : UNESCO, 80 p., doi: 10.54675/JKJB9835.
11. Aleshkovskiy, I.A., Gasparishvili, A.T., Narbut, N.P., Krukhmaleva, O.V., Savina, N.E. (2024). Russian Students on the Opportunities and Limitations of Artificial Intelligence in Learning. *Vestnik RUDN. Seriya: Sotsiologiya = RUDN Journal of Sociology*. Vol. 24, no. 2, pp. 335-353, doi: 10.22363/2313-2272-2024-24-2-335-353 (In Russ., abstract in Eng.).
12. Buyakova, K.I., Dmitriev, Ya.A., Ivanova, A.S., Feshchenko, A.V., Yakovleva, K.I. (2024). Attitudes of Students and Teachers Toward the Use of Artificial Intelligence Tools in a University]. *Obrazovanie i nauka = The Education and Science Journal*. Vol. 26, no. 7, pp. 160-193, doi: 10.17853/1994-5639-2024-7-160-193 (In Russ., abstract in Eng.).
13. Koshkina, E.A., Bordovskaya, N.V., Gnedykh, D.S., Khromova, M.A., Dem'yanchuk, R.V., Iskhakova, M.P., Balyshv, P.A. (2025). Generative Artificial Intelligence in Higher Education: A Review of theoretical Approaches and Application Practices. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 36-57, doi: 10.31992/0869-3617-2025-34-6-36-57 (In Russ., abstract in Eng.).
14. Kuz'minov, Ya.I., Kruchinskaya, E.V., Gruzdev, I.A., Naumov, A.A. (2025). Laggards and Front-runners: How Students Use Generative Artificial Intelligence for Educational Purposes. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 9-35, doi: 10.31992/0869-3617-2025-34-6-9-35 (In Russ., abstract in Eng.).
15. Tikhonova, N.V., Pomortseva, N.P. (2025). Final Qualifying Work at a University in the Context of the Spread of Artificial Intelligence: Students' Perspective. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 112-135, doi: 10.31992/0869-3617-2025-34-6-112-135 (In Russ., abstract in Eng.).
16. Zemtsov, D.I., Gruzdev, I.A. (2025). "Digital Centaur": Human-AI Collaborative Learning at the University. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 10, pp. 47-62, doi: 10.31992/0869-3617-2025-34-10-47-62 (In Russ., abstract in Eng.).
17. Oncioiu, I., Bularca, A.R. (2025). Artificial Intelligence Governance in Higher Education: The Role of Knowledge-Based Strategies in Fostering Legal Awareness and Ethical Artificial Intelligence Literacy. *Societies*. Vol. 15, no. 6, article no. 144, doi: 10.3390/soc15060144.
18. Hackl, V., Müller, A.E., Sailer, M. (2026). The AI Literacy Heptagon: A Structured Approach to AI Literacy in Higher Education. *Computers and Education: Artificial Intelligence*. Vol. 10, article no. 100540, doi: 10.1016/j.caeai.2026.100540.
19. Zhang, C., Magerko, B. (2025). Generative AI Literacy: A Comprehensive Framework for Literacy and Responsible Use. *arXiv:2504.19038*, doi: 10.48550/arXiv.2504.19038.
20. Zamfirescu-Pereira, J.D., Wong, R.Y., Hartmann, B., Yang, Q. (2023). Why Johnny Can't Prompt: How Non-AI Experts Try (and Fail) to Design LLM Prompts. In: *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. New York : Association for Computing Machinery, article no. 437, pp. 1-21, doi: 10.1145/3544548.3581388.

21. Federiakin, D., Molerov, D., Zlatkin-Troitschanskaia, O., Maur, A. (2024). Prompt Engineering as a New 21st Century Skill. *Frontiers in Education*. Vol. 9, article no. 1366434, doi: 10.3389/educ.2024.1366434.
22. Carrasco-Sáez, J.L., Contreras-Saavedra, C., San-Martín-Quiroga, S., Contreras-Saavedra, C.E., Viveros-Muñoz, R. (2025). Analyzing Higher Education Students' Prompting Techniques and their Impact on ChatGPT's Performance: An Exploratory Study in Spanish. *Applied Sciences*. Vol. 15, no. 14, article no. 7651, doi: 10.3390/app15147651.
23. Qu, Y., Tan, M.X.Y., Wang, J. (2024). Disciplinary Differences in Undergraduate Students' Engagement with Generative Artificial Intelligence. *Smart Learning Environments*. Vol. 11, article no. 51, doi: 10.1186/s40561-024-00341-6.
24. Ferdman, A., Ratti, E. (2024). What Do We Teach to Engineering Students: Embedded Ethics, Morality, and Politics. *Science and Engineering Ethics*. Vol. 30, article no. 7, doi: 10.1007/s11948-024-00469-1.
25. Nass, C., Moon, Y. (2000). Machines and Mindlessness: Social Responses to Computers. *Journal of Social Issues*. Vol. 56, no. 1, pp. 81-103, doi: 10.1111/0022-4537.00153.
26. Epley, N., Waytz, A., Cacioppo, J.T. (2007). On Seeing Human: A Three-Factor Theory of Anthropomorphism. *Psychological Review*. Vol. 114, no. 4, pp. 864-886, doi: 10.1037/0033-295X.114.4.864.
27. Gong, L. (2008). How Social Is Social Responses to Computers? The Function of the Degree of Anthropomorphism in Computer Representations. *Computers in Human Behavior*. Vol. 24, no. 4, pp. 1494-1509, doi: 10.1016/j.chb.2007.05.007.
28. Waytz, A., Heafner, J., Epley, N. (2014). The Mind in the Machine: Anthropomorphism Increases Trust in an Autonomous Vehicle. *Journal of Experimental Social Psychology*. Vol. 52, pp. 113-117, doi: 10.1016/j.jesp.2014.01.005.
29. Gambino, A., Fox, J., Ratan, R.A. (2020). Building a Stronger CASA: Extending the Computers Are Social Actors Paradigm. *Human-Machine Communication*. Vol. 1, pp. 71-86, doi: 10.30658/hmc.1.5.
30. Yankouskaya, A., Babiker, A.B., Rizvi, S.W.F., Alshakhsi, S., Liebherr, M., Ali, R. (2025). LLM-D12: A Dual-Dimensional Scale of Instrumental and Relational Dependencies on Large Language Models. *ACM Transactions on the Web*. Advance online publication, doi: 10.1145/3765895.
31. Pikhart, M., Al-Obaydi, L.H. (2025). Reporting the Potential Risk of Using AI in Higher Education: Subjective Perspectives of Educators. *Computers in Human Behavior Reports*. Vol. 18, article no. 100693, doi: 10.1016/j.chbr.2025.100693.
32. Huang, D., Hash, N., Cummings, J.J., Prena, K. (2025). Academic Cheating with Generative AI: Exploring a Moral Extension of the Theory of Planned Behavior. *Computers and Education: Artificial Intelligence*. Vol. 8, article no. 100424, doi: 10.1016/j.caeai.2025.100424.
33. Pangestu, F.N.N., Hamidah, I., Widaningsih, L., Abdullah, A.G., Saputra, N.A. (2026). A Scoping Literature Review of Prompt Engineering for Bridging Students' AI Literacy in Higher Education. *Computers and Education: Artificial Intelligence*. Vol. 10, article no. 100581, doi: 10.1016/j.caeai.2026.100581.
34. Ivanova, A.E., Tarasova, K.V., Talov, D.P. (2025). Between Interest and Skill: How Students Perceive and Use AI. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 8-9, pp. 9-32, doi: 10.31992/0869-3617-2025-34-8-9-9-32 (In Russ., abstract in Eng.).
35. Consoli, T., Petko, D. (2025). Which Educational Approaches Predict Students' Generative AI Confidence and Responsibility? *Computers and Education: Artificial Intelligence*. Vol. 9, article no. 100431, doi: 10.1016/j.caeai.2025.100431.

36. Schutte, N.S., Li, H. (2025). The Role of Self-Efficacy and Curiosity in Student Use of Artificial Intelligence (AI). *International Journal of Educational Technology in Higher Education*. Vol. 22, article no. 73, doi: 10.1186/s41239-025-00574-6.
37. Jongkind, R., Bikker, Y., Meinema, J., Broens, T. (2025). Studying with GenAI: Cross-Sectional Study on Usage Patterns, Needs, Competencies, and Ethical Perspectives of Medical Informatics Students. *Frontiers in Education*. Vol. 10, article no. 1658415, doi: 10.3389/feduc.2025.1658415.
38. Hadinejad, N., Sperling, K., McGrath, C. (2025). Generative AI Chatbots in Higher Education: Student Experiences and Perceived Ethical Challenges. *Computers and Education Open*. Vol. 9, article no. 100311, doi: 10.1016/j.caeo.2025.100311.

Acknowledgement. The study was implemented in the framework of the Basic Research Program at HSE University in 2026 (HSE-BR-2025-009).

Статья поступила в редакцию 29.04.2026

Принята к публикации 04.06.2026

Сведения для авторов

К публикации принимаются статьи, как правило, не превышающие 40000 знаков.

Название файла со статьёй – фамилии и инициалы авторов. Таблицы, схемы и графики должны быть представлены в формате MS Word (с возможностью редактирования) и вставлены в текст статьи. Подписи к рисункам, графикам, диаграммам, таблицам должны быть продублированы на английском языке.

Рукопись должна включать следующую информацию *на русском и английском языках*:

- название статьи (не более шести-семи слов);
- сведения об авторах (ФИО полностью, учёное звание, учёная степень, должность, ORCID, Researcher ID, e-mail, название организации с указанием полного адреса и индекса);
- аннотация и ключевые слова (отразить цель работы, методы, основные результаты и выводы, объём – не менее 250–300 слов, или 20–25 строк); весь блок на английском языке должен быть прочитан и одобрен специалистом-лингвистом или носителем языка;
- литература (15–25 и более источников). Ссылки даются в порядке упоминания.

В целях расширения читательской аудитории и выхода в международное научно-образовательное пространство рекомендуется включать в список литературы (References) зарубежные источники. Важно: при оформлении References имена авторов должны указываться в оригинальной транскрипции (не транслитом!), а название источника – в том виде, в каком он был опубликован. Если источник имеет DOI, его следует указывать.

Если в статье имеется раздел «Благодарности» (Acknowledgement), то в англоязычной части статьи следует разместить его перевод на английский язык.

Рекомендуем перед отправкой рукописи в редакцию убедиться, что статья оформлена по нашим правилам.

Факторы интеграции искусственного интеллекта в исследование: кейсы российских лабораторий

Научная статья

DOI: 10.31992/0869-3617-2026-35-6-149-168

Сазонов Артём Алексеевич – младший научный сотрудник Научно-исследовательской лаборатории археологии и этнографии Сибири, ORCID: 0009-0004-6088-057X, artemiy_sazonov@mail.ru

Попова Евгения Владимировна – канд. полит. наук, доцент, заведующая кафедрой антропологии и этнологии Факультета исторических и политических наук, ORCID: 0000-0003-1716-8849, iam.e.popova@yandex.ru

Мацепуро Дарья Михайловна – старший научный сотрудник Научно-исследовательской лаборатории археологии и этнографии Сибири, ORCID: 0000-0002-9809-082X, daria.matsepuro@mail.tsu.ru

Томский государственный университет, Томск, Россия

Адрес: 634050, г. Томск, ул. Ленина, д. 36

***Аннотация.** Для понимания факторов успешной интеграции технологий искусственного интеллекта (ИИ) как исследовательского метода в деятельность лаборатории необходимо обратиться к существующему опыту научных коллективов, проанализировать ключевые аспекты их работы и опыт внедрения искусственного интеллекта, успешный или нет. Цель настоящего исследования состоит в том, чтобы выявить и систематизировать ключевые переменные, определяющие возможность включения искусственного интеллекта в работу научных коллективов. Для этого мы используем подход множественного кейс-стади. В центре внимания данного исследования четыре исследовательских коллектива сибирских научных и образовательных организаций. Научные группы работают или работали над проектами в разных предметных областях: физике, истории, медицине.*

Мы выявили три взаимосвязанных условия успешного внедрения искусственного интеллекта. Во-первых, возможность гармонизации исследовательской задачи и формата данных – исследователи начинали активно интегрировать ИИ в свою работу, когда традиционные методы становились технически неэффективными. Во-вторых, формирование коллаборации со специалистами в области работы с данными и программистами, которая во многом зависела от активной организационной политики, позволяя приобрести необходимые компетенции. И, как следствие, двусторонняя коммуникация между исследователями и специалистами по ИИ позволила закрепить опыт применения технологии для действующего и будущих проектов. Таким образом, внедрение ИИ оказывается не столько линейным, технологически обусловленным процессом, сколько согласованием исследова-

тельской задачи, имеющихся ресурсов, компетенций, коммуникативных навыков, организационной культуры и дисциплинарных норм валидации полученного знания.

Ключевые слова: научная инфраструктура, искусственный интеллект, социальные исследования лабораторий, STS, стандарты научного знания

Для цитирования: Сазонов А.А., Попова Е.В., Мацепуро Д.М. Факторы интеграции искусственного интеллекта в исследование: кейсы российских лабораторий // Высшее образование в России. – 2026. – Т. 35, № 6. – С. 149–168. – DOI: 10.31992/0869-3617-2026-35-6-149-168.

For Factors of AI Integration in Research: Case Study of Russian Laboratories

Original article

DOI: 10.31992/0869-3617-2026-35-6-149-168

Artem A. Sazonov – Junior Researcher, Laboratory of Archaeology and Ethnography of Siberia, ORCID: 0009-0004-6088-057X, artemiy_sazonov@mail.ru

Evgeniya V. Popova – Cand.Sci. (Political Sciences), Associate Professor, Head of the Department of Anthropology and Ethnology, Faculty of Historical and Political Sciences, ORCID: 0000-0003-1716-8849, iam.e.popova@yandex.ru

Daria M. Matsepuro – Cand.Sci. (History), Senior Researcher, Laboratory of Archaeology and Ethnography of Siberia, ORCID: 0000-0002-9809-082X, Researcher ID: AAF-6266-2019, daria.matsepuro@mail.tsu.ru

National Research Tomsk State University, Tomsk, Russian Federation

Address: 36 Lenin ave., Tomsk, 634050, Russian Federation

Abstract. To understand the factors of successful integration of artificial intelligence into research process, it is necessary to examine the existing experience of scientific collectives, analyze key aspects of their work and the experience of implementing artificial intelligence (AI), whether successful or not. The purpose of this study is to identify the key factors that determine the success of artificial intelligence integration into the research activities of scientific collectives.

To achieve the research objectives, we employ a multiple case study approach and conduct a series of semi-structured in-depth interviews. The focus of this study is four research collectives from Siberian scientific institutions. The scientific groups have been working or have worked on projects from different disciplines: physics, history, and medicine.

We identified three interconnected conditions for successful artificial intelligence implementation. First, the possibility of harmonizing the research task and data format. Active integration of AI into the work of research collectives started when traditional methods had become inefficient from technological standpoint. Second, active organizational policy allowed research collectives to establish collaborations with programmers, acquiring necessary competencies. Third, two-way communication between researchers and AI specialists enabled consolidation of experience in applying the technology, reducing costs, and integrating AI specialists into the scientific group. Thus, the implementation of AI proves to be not so much a linear, technologically determined process, but rather a coordination of the research task, available resources, competencies, organizational culture, and disciplinary norms for validating acquired knowledge.

Keywords: *research* infrastructure, artificial intelligence, laboratory social studies, STS, standards of scientific knowledge

Cite as: Sazonov, A.A., Popova, E.V., Matsepuro, D.M. (2026). Factors of AI Integration in Research: Case Study of Russian Laboratories. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 35, no. 6, pp. 149-168, doi: 10.31992/0869-3617-2026-35-6-149-168 (In Russ., abstract in Eng.).

Введение

Искусственный интеллект (ИИ) – один из ключевых инструментов современной науки и промышленности – меняет анализ данных и моделирование сложных систем. Одновременно растёт скепсис учёных к ИИ-решениям. Они представляют «чёрный ящик»: работа алгоритмов ИИ непрозрачна, это подрывает доверие и тормозит внедрение перспективных методов. Алгоритм в процессе самообучения меняется, суть этих изменений мало кто понимает [1].

США и Китай инвестируют в научный ИИ. Более 35% публикаций по изучению влияния ИИ на общество были выпущены в этих странах [2]. Понимание факторов успешной интеграции технологий – вопрос стратегического преимущества. При этом внедрение ИИ неравномерно: одни лаборатории активно пользуются алгоритмами ИИ, другие не включают их в работу.

Стремительное развитие технологии обнажает системную проблему. Обнаруживается разрыв между возможностями алгоритмов и готовностью исследователей их использовать, между технологическим давлением и организационными и иными барьерами. Наблюдается неравномерность внедрения ИИ: географическая, дисциплинарная, организационная. Её причины неочевидны и не сводятся к наличию данных или вычислительных мощностей; предположительно, они кроются в малоизученных механизмах принятия решений на уровне научных групп. В фокусе данной работы находится вопрос: по каким критериям различные научные коллективы принимают решение интегрировать ИИ в исследования, и что определяет успех или неудачу внедрения. Кажется, это зависит от стратегии исследования, накопления и перераспределения имеющихся ресурсов, а также

от культурных и социальных контекстов деятельности учёных.

Методологически мы исходим из базового представления социальных исследований науки и технологий (STS), что деятельность исследователя определяется не внутренней логикой науки, но положением в транснаучных полях – сетях отношений, которые выходят за пределы научного сообщества [3], или культурными паттернами, которые определяют способы производства знания в разных дисциплинах. Иначе знание производится в повседневных практиках, и новые методы разными способами встраиваются в них. Решение учёных не сводится здесь к индивидуальной рациональности, но связано со сложным переплетением социальных, культурных, дисциплинарных норм и материальных условий. Расположение устройств в лаборатории связано не только с исследовательской необходимостью, но и с переговорами с другими сотрудниками университета; дизайном исследования, в большей степени отвечающего интересам грантодателей и их представлениям о стандартах научного знания. Общая рамка исследования – STS позволяет сфокусироваться не только на технических вопросах (точность, доступность, виды данных) или на социальных переменных (власть, институты, ресурсы), но на том, как в решении задействуется эпистемическая идентичность учёного. В принятии решения может быть важным, что считается «хорошим» знанием и правильными научными практиками в данной дисциплине [3]; каковы локальные культуры (нормы, иерархии) в контексте локальной научной группы и российской институциональной структуры в целом [4]; и наконец, как материальная организация исследовательского пространства опосредует интеграцию ИИ?

Возможности и риски применения ИИ в работе исследователя: обзор литературы

Сейчас выработан обширный корпус литературы, который описывает способы применения ИИ в академической среде. Много внимания привлекают вопросы этичности использования генеративного ИИ в образовании, его потенциал для тьюторского сопровождения, оценивания студентов [5–7]. В контексте научных исследований в основном литература посвящена возможностям, которые открывает ИИ перед учёными в процессе исследования [8–11]. Их можно обобщить по трём категориям: организация рабочего процесса, анализ данных и производство публикаций – создание текстов, инфографики, перевод¹. Первая группа показывает позитивную динамику по числу работ, связано это с ростом числа ИИ, специализированных для решения научных задач [12]. Фокус здесь находится прежде всего на новых инструментах, а не осмыслении процесса и результата их включения в работу исследователя. Публикуются модели, разработанные как для работы в узкодисциплинарных рамках: медицине [13; 14], физике [15], социальных исследованиях [16], так и в научных исследованиях в целом [17]. Отмечается потенциал ИИ в поиске и анализе научной литературы. Предлагаются методические решения для обработки массивов публикаций [18] и формулирования гипотез². Например, исследователи отмечают его высокую эффективность для анализа литературы [19; 20]. Но есть и критика такого подхода: уменьшение проверяемости и надёжности исследования [21].

Также предлагается инфраструктурный подход к внедрению ИИ, например для создания и обработки баз данных [22]. Применение систем ИИ для анализа научных данных включает в себя отбор и классификацию

эмпирического материала. В социальных исследованиях предлагается использовать большие лингвистические модели, чтобы быстрее формировать темы и облегчать качественное кодирование [23; 24] или автоматический анализ интервью [25; 26]. С. Юнцзюнь и коллеги приводят обзор возможностей алгоритмов в разных областях науки. Например, в медицине ИИ применяется для разработки лекарств, прогнозирования течения заболеваний, а в материаловедении для симуляции поведения новых материалов в различных условиях [9].

Инструменты ИИ могут вычитывать академические тексты, корректировать их, составлять аннотации [27]. Некоторые исследователи также используют нейросети для написания грантовых заявок [27]. ИИ автоматизирует визуализацию данных: создаёт диаграммы, графики на основе необработанных данных [27]. Нейросетевые переводчики облегчают и ускоряют коммуникацию между учёными, не владеющими языками международного общения [28].

Исследователи выделяют риски, связанные с внедрением ИИ в научный процесс. Часть из них касается этических вопросов. Например, системы машинного зрения могут укреплять существующие в обществе расовые и гендерные стереотипы [29]. Самые актуальные риски связаны с генеративными моделями ИИ. Они могут влиять на разнообразии исследовательских культур. Например, двумя способами: предлагая методы исследования, которые больше подходят для применения искусственного интеллекта; и отбирая идеи, которые ИИ способен выражать [30]. Например, широкое распространение больших лингвистических моделей приводит к тому, что в публикациях чаще встречаются специфичные для генеративно-

¹ AI Insights 2024. [Электронный ресурс]. Elsevier. 2025. URL: https://assets.ctfassets.net/o78em1y1w4i4/6BWRibYJNQLYkKWwKw7SVf/64c04b53ca9cc0795ac811f583f7eebb/Insights_2024_Attitudes_To_AI_Full_Report.pdf (дата обращения: 08.12.25)

² Accelerating scientific breakthroughs with an AI co-scientist [Электронный ресурс]. Google Research. 2025. URL: <https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/> (дата обращения: 08.12.25)

го ИИ формулировки и слова [31], а включение их в исследовательские инфраструктуры потенциально ведёт к потере возможности использовать альтернативные инструменты [32]. Генеративный ИИ способен понимать и применять естественный язык, что ведёт к снижению уровня когнитивных затрат при использовании таких инструментов, что создаёт риск увеличения числа статей с низким научным уровнем [33].

Среди исследователей нет согласия о возможностях, которые предоставляет ИИ, и связанных с ним рисков. Его можно достичь, фокусируясь на том, как учёные принимают решения о внедрении ИИ в исследования. Литературы, посвящённой этому вопросу, мало. Редкие исследователи рассматривают институциональные и социальные факторы, влияющие на решение использовать ИИ. Д. Лю и Х.В. Джагадиш предполагают, что на принятие исследователями технологии ИИ могут повлиять вычислительные мощности и компетенции исследователей [34]. С. Бьянкини и коллеги соглашались с необходимостью создания технических инфраструктур, но не считают их решающим фактором. Исследователи подчёркивают важность институционального контекста: как много исследовательских коллективов уже используют ИИ и сколько молодых учёных работают в составе научной организации [35].

В литературе нет консенсуса о факторах включения ИИ в исследовательскую практику, тема остаётся малоизученной. Также в литературе присутствует фокус на системах генеративного ИИ и больших лингвистических моделях. Разработки в этой сфере флагманские, однако из виду теряются иные системы ИИ, такие как компьютерное зрение. Генеративный ИИ представлен больше других технологий в социальных исследованиях по этой теме [36].

В приведённых выше исследованиях ИИ рассматривается как сложившаяся технологическая данность. Его влияние на челове-

ское общество не может быть тотальным. В результате существует риск получить корпус материалов, который основан на неполных данных, и потерять из виду различные стадии интеграции технологии.

Российское поле со своей академической и этической спецификой – более мягким отношением к этическому регулированию, опорой на рекомендательные, а не регуляторные нормы в области ИИ – остаётся слабо изученным [37].

Таким образом, цель исследования – выявить и систематизировать ключевые факторы, определяющие успешность интеграции ИИ в повседневную исследовательскую деятельность научных коллективов. Для этого мы представляем существующие практики интеграции ИИ у изученных нами научных коллективов. Затем выявляем ключевые факторы, влияющие на этот процесс, и наконец, представляем барьеры, приводящие к неудаче или осложнению интеграции ИИ в исследования.

Методы сбора и анализа данных

В исследовании используется подход множественного кейс-стади, что позволяет оценить процесс с позиций разных исследовательских групп. В центре внимания четыре исследовательских коллектива сибирских научных организаций из разных предметных областей: историки, медицинские биофизики, теплофизики и физики-фотонщики. Они представляют разные структуры: университетские научные группы и структуры и лаборатории РАН³. Выборка лабораторий была стихийной, поскольку число научных коллективов, имеющих опыт применения ИИ в проведении исследований на момент исследования, было невелико. Лаборатории искали с помощью университетских Центров ИИ, научных управлений и по личным контактам исследователей.

Мы собирали данные с помощью полуструктурированных интервью. Гайд состоял

³ Мы используем здесь термины «лаборатория» и «научная группа» как синонимы.

из нескольких блоков: структура лаборатории и позиция информанта в ней; представления об ИИ и история его интеграции в исследовательскую работу; этические импликации использования ИИ. Всего проведено 12 интервью, длительностью от 40 минут до полутора часов. Число интервью для каждой лаборатории невелико. Обычно это руководитель или научный сотрудник – идеолог включения ИИ в исследование, аспирант, который работает с ИИ, потому что ему «дали тему», или ИИ-специалист. В случае историков был разговор с двумя сотрудниками научной группы.

Материалы интервью отличаются от кейса к кейсу, поэтому было проведено качественное кодирование данных. Мы выявили первичные коды во всех интервью: научная задача, данные, компетенции, коммуникации, результаты чего представлены в разделе «Эмпирическая часть». На этапе анализа кейсов мы работали с первичными кодами вокруг теоретических категорий эпистемических культур, типов коллабораций разнообразных участников и ситуативной рациональности.

При выборе стратегии исследования научные группы опираются на тип и особенности данных, которые им предстоит анализировать, например, как показывают И. Стайнхарт и коллеги, практики учёных по обработке открытых данных могут варьироваться от их типа [38]. Зависимость от типа и качества данных особенно характерна в исследованиях с использованием ИИ, которые опираются на статистические методы. Такие вопросы решаются на этапе дизайна исследования: какие методы будут использоваться и к каким данным применяться [39]. Соответственно, в центре внимания оказываются три взаимосвязанных аспекта, вокруг которых выстраивается наше описание работы лабораторий с ИИ как методом исследования: характер и пригодность данных для обработки алгоритмами, наличие компетенций, а также выстраивание эффективной коммуникации между учёными-предметниками и специалистами по ИИ. Кодирова-

ние происходило в соответствии с описанной схемой. Мы использовали программу *QualCoder*, чтобы структурировать эмпирический материал, организовать его вокруг трёх центральных категорий, описывающих факторы интеграции искусственного интеллекта в исследование.

В начале каждого интервью мы предупреждали информантов о рисках. Главный риск в этих условиях – деанонимизация. Согласие на участие в исследовании и публикацию цитат фиксировалось на диктофон. Для минимизации рисков мы указываем минимальное количество информации о респондентах – только сферу научной специализации.

Эмпирическая часть:

четыре кейса внедрения ИИ

Лаборатория 1. Историки, университет. Группа историков занимается исследованием медицины традиционными методами с разнородными источниками. До описываемого проекта они не применяли цифровые методы для анализа источников. В рамках проекта они изучали корпус газет, используя алгоритмы распознавания текста и математические методы анализа, чтобы определять, что или кого авторы связывали со вспышками чумы и какова связь уровня радикализации общества с эпидемической ситуацией. Историки реализовывали проект в коллаборации с центром ИИ, открытым в университете. Историки использовали ИИ для анализа обширного корпуса оцифрованных газет. Целью было обнаружение скрытых паттернов в текстах, а именно выделение понятий, обозначающих разные подходы к эпидемии чумы: «разных понятий, с одной стороны, обозначающих жёсткий и ксенофобный подход к ответу на эпидемию [чумы], и, с другой стороны, мягкий подход» (историк). Задача была переведена в формализуемую плоскость, что позволило подключить вычислительные методы.

Руководство университета финансово стимулировало проекты с ИИ, поэтому

группа решила использовать нетипичный метод. Изначально коллектив относился к этой идее с опасениями: *«Мы [историки], разумеется, скептически к этому относимся, ну что можно с этим сделать, как...»*. Чтобы начать работу один из проректоров организовал встречу учёных с датасайентистами, которые *«очень динамично, энергично подошли, и в ответ на то, что я рассказал о нашем проекте, начали что-то предлагать..., и что мы можем ... получить»* (руководитель группы историков). Это позволило сформулировать пригодную для применения ИИ научную задачу в сфере компетенций историков и уменьшить скепсис в отношении «технарей».

Взаимодействие с ИИ-специалистами строилось по модели ясного разделения труда. Учёные-предметники формулировали задачу, размечали тексты для обучения программы и интерпретировали результаты, а датасайентисты *«предложили механизмы, алгоритмы, программы»* для решения проблемы. Но, как подчёркивает руководитель историков, *«все задачи, всю творческую составляющую формулировали мы»*, формируя чёткую границу между наукой и вспомогательными инструментами, закрепляя за человеком креативную и интерпретативную функции.

Другой сложностью стало качество данных. Несмотря на обширный оцифрованный корпус дореволюционных газет его машиночитаемость оказалась низкой: *«... Практически все газеты на данный момент оцифрованы. Многие из них оцифрованы с распознаванием текста. Но ...программы для распознавания текстов оцифрованных газет очень несовершенны»* (историк). Это привело к тому, что *«используя автопоиск, невозможно обнаружить слово, даже если оно есть»*. ИИ-специалистам пришлось обрабатывать сырые данные. Из-за небольшого финансирования проекта продвинутые методы оказались недоступны, «технари» использовали альтернативные современным ИИ-инструментам методы.

Руководитель группы эмоционально говорит о перспективах применения ИИ, но тут же запал снижается, так как это дорого, и финансирования проектов недостаточно для амбициозных задач. Историки опасаются, что научное сообщество в России не примет их работы. Информант-историк отмечает принципиальное отличие машинной обработки от человеческой: *«Человек есть человек, машина есть машина. Если человек последовательно десятки и сотни номеров газет просматривает, он будет ошибаться. ...Машина не ошибётся, если алгоритмы работают хорошо»*. Руководитель коллектива видит задачу-максимум в том, чтобы *«добиться усиления потенциала исторического источника, чтобы он давал качество, которое невозможно получить без технических средств. Историки прошлого, настоящего уже добились максимума. Как выйти за эту границу?»*.

После завершения проекта применение ИИ в лаборатории историков прекратилось. Технология осталась на аутсорсе, и глубокой трансформации исследовательской инфраструктуры не произошло.

Лаборатория 2. Медицинские биофизики, лаборатория РАН. Коллектив пытался определять инфаркт на основе радиомических показаний. Радиомика использует характеристики градаций цвета серошальных изображений для предсказания свойств изображённого объекта. Количество параметров при анализе КТ-снимков, с учётом количественной выборки, может достигать, по словам руководителя лаборатории, от 93 до десятков тысяч. Обработка такого массива данных привычными статистическими методами слишком трудоёмка. Исследования на таком объёме данных сейчас проводятся с помощью технологий нейросетей и машинного обучения [40]: *«Данные вносят в нейросеть, которая отбрасывает лишние показатели. Либо объединяет в группы...»* (биофизик 3). Более того с появлением методов визуального распознавания цвета на снимках с помощью ИИ появляются бес-

прецедентные возможности для обработки изображений, а не только моделирования таблиц. Подход активно используется в онкологии, но слабо представлен в других исследованиях. В результате проекта была защищена диссертация и опубликовано 4 статьи в журналах Q4–Q3. Все участники знают о возможностях ИИ в радиомике, с энтузиазмом рассказывали о них, о том, что у них было много надежд, связанных с включением ИИ, но в результате так эти методы и не использовали, на что повлияло несколько факторов.

Исследовательский коллектив не обладал компетенциями в разработке ИИ: «... по мере углубления и получения дополнительных знаний, мы просто все за голову схватились. Мы так не умеем, нам нужны программисты...». В лаборатории были специалисты биокибернетики, но не было никого, кто бы имел экспертизу в разработке нейросетей. Возможности нанять специалистов в области ИИ не было из-за отсутствия финансирования. В отличие от других групп медики в регионе не имеют возможности собрать большой комплекс ретро-данных. Институт ввёл цифровую историю болезней относительно недавно, результаты прошлых клинических обследований не были оцифрованы. Оцифровка для ИИ *«очень трудозатратна. Кто это будет делать? ... Там нужно тысячи, десятки тысяч [историй болезни]»* (руководитель). Сейчас система здравоохранения не может предоставить долговременные данные пациентов в объёме, чтобы полноценно применять к ним ИИ-анализ. Ожидаемые результаты алгоритма радиомики вызывают опасения: *«Все прогностические модели, которые строились на основании данных [радиомики], которые имели статистическую достоверность, это какая-то игра цифр»* (биофизик 2). Исследователь признавался, что вычисления привычных устройств и программного обеспечения такой же «чёрный ящик», но он понимает физические принципы, которые за ними

стоят. Радиомические показатели, хоть и имели статистическую значимость, но не были связаны ни с какими известными ему радиологическими процессами.

Важную роль играет система доступа к данным. Помимо получения информированного согласия в исследованиях, связанных с человеком, исследователь должен доказать необходимость сбора данных клиницисту, который может спросить: *«...а зачем? Всё уже давно исследовано, всё известно»* (биофизик 2). Биофизики в медицинских исследованиях не могут собирать клинические данные и интерпретировать их без участия врача.

В результате ИИ как исследовательский метод не был применён ни в каком виде.

Лаборатория 3. Теплофизическая лаборатория, университет. Занимается экспериментальными работами в сфере предикции пожаров и фундаментальных исследований процессов горения разных материалов. С помощью ИИ разработана система предупреждения возгорания, на основе анализа данных прошедших экспериментов и вновь поступающих данных с датчиков разных участков исследования.

Стимулом для внедрения ИИ в научную практику стал характер исследований, связанный с анализом множества параметров в реальном времени: *«Когда доходишь до определённого уровня с большим количеством данных, начинаешь задумываться, как оперативно их обрабатывать...»* (руководитель). Это подтолкнуло коллектив к поиску инструментов для работы с большими данными. Изначально работа с ИИ не планировалась. Датасайентисты должны были помочь настроить статистические таблицы. Затем появилась идея автоматизации, а уже после попыток сделать данные читаемыми, стало понятно, что необходима работа с ИИ [подробнее о кейсе: 1]. После успешного опыта применения нового инструмента технология была распространена на три новых проекта в разных междисциплинарных областях, включая медицину.

В этом кейсе самый глубокий уровень коллаборации между учёными-предметниками и ИИ-специалистами. Изначально взаимодействие было сложным: на знакомство последних с сутью эксперимента ушло около полугода. Ключевым элементом стали регулярные встречи представителей групп, на которых вырабатывался общий язык и выявлялись места пересечения экспертиз. Со временем коммуникация наладилась настолько, что один из ИИ-специалистов поступил в аспирантуру по физике.

Главной проблемой, как и у историков, оказалось качество данных. Данные экспериментов, хранившиеся во множестве файлов *Excel*, оказались несопоставимыми: «...видно, что [таблицы] все разные по структуре. Где-то 10 столбцов с данными, а где-то 20. Хотя бокс, где проводились эксперименты, один» (ИИ-специалист). Решение нашлось в радикальном изменении подхода: датасайентисты предложили отказаться от *Excel*-таблиц учёных, взять данные напрямую из контроллера: «[оказалось], эти *Excel* получались из баз данных контроллера. И мы говорим, ребята, зачем мы голову морочили с этими *Excel*, давайте мы просто будем брать этот источник» (ИИ-специалист). Это нарушило привычные учёным формы работы, но открыло путь для ИИ.

Пока что это единственный случай, в котором успешный опыт привёл к трансформации лаборатории. Коллектив начал обустривать инфраструктуру работы ИИ, обучать аспирантов одновременно ИИ и физике. По мере реализации проектов интеграция ИИ стала требовать меньше когнитивных усилий, времени и финансов. Изменился кадровый состав и формат работы: началась подготовка студентов на стыке естественно-научных проектов и ИИ.

Лаборатория 4. Физики-лазерщики, университет. В этом кейсе ИИ изначально рассматривался как основной метод исследования. Руководитель лаборатории, разрабатывающей зеркала для твердотельных лазеров, обратился к руководителю уни-

верситетского центра ИИ с предложением о сотрудничестве.

Задача состояла в разработке многослойного покрытия (фотонного кристалла) для зеркал, отражающих свет в нужном диапазоне. Поиск оптимальной структуры покрытия – сложная вычислительная задача, требующая перебора множества комбинаций. Руководитель уверен, что ИИ ускорит процессы расчёта и поиска моделей. Совместно с центром ИИ решили разработать «эволюционный алгоритм», имитирующий эволюционные процессы, который должен был подсказать наиболее удачное решение. Как и для других новаций в науке в разных странах, руководитель поручил задачу своему аспиранту, который работает плотно с ИИ-специалистом из Центра ИИ вуза. Руководитель еженедельно слушает отчёт аспиранта и включается в проект и в коммуникацию лишь в случае сложных вопросов. Таким образом, коммуникация строится на паритетных началах со-разработки. Учёный-предметник ставит задачи, а ИИ-специалист и аспирант совместно разрабатывают новый инструмент. Проект ещё не завершён, результаты этого кейса ещё не получены, но стороны настроены оптимистично. Никто из участников не смог обозначить какие-либо явные организационные сложности работы. В случае успеха ИИ станет не просто инструментом для решения конкретной задачи, а новым методом проектирования, встроенным в исследовательский процесс. Авторы смогли определить ряд познавательных задач, которые возможно передать алгоритму, сообщив ему роль экспериментатора.

Ключевые условия и барьеры внедрения ИИ: результаты сравнения

Обобщая опыт четырёх описанных лабораторий, можно выделить ряд факторов, которые определяют успешность и характер интеграции искусственного интеллекта в исследовательскую практику. Есть две категории: условия, способствующие внедрению, и барьеры, его блокирующие.

Позитивные факторы внедрения

Гармонизация задач и форматов данных

Внедрение систем ИИ в науку связано с набором условий, которые учёные вынуждены учитывать, принимая решение об интеграции этих технологий. Важно, что условия дисциплинарно специфичны, отражая повседневность эпистемических культур. Прежде всего, в основе решения лежит научная проблема. Учёные рассматривают ИИ не как цель, а как возможное решение для определённого типа задач. ИИ демонстрирует высокую эффективность в чётко сформулированных, повторяемых операциях: поиск, классификация, кластеризация, прогнозирование. ИИ применяется, когда задача, часто связанная с прогнозированием, предполагает обработку информации в объёмах, недостижимых для человеческого восприятия. Это может быть поиск скрытых паттернов в больших массивах текстовых корпусов, как у историков, или анализ множества параметров в реальном времени, как в случае с физиками и медиками. Таким образом, процессы датафикации исследований становятся стимулом для работы с новым, проблемным инструментом.

Как показал опыт изученных коллективов, научная проблема должна быть формализуемой, чтобы использовать ИИ. В случае историков – выделение полярных понятий в текстах, у фотонщиков – оптимизация структуры фотонного кристалла, у теплофизиков – прогнозирование возгораний на основе оценки поведения горящих веществ. Большое количество функций сохраняются за человеком. У историков это разметка данных и их интерпретация, у физиков – работа по унификации собираемых в эксперименте данных, пригодных для анализа методами ИИ. Теплофизикам для обучения модели пришлось выстраивать коммуникацию с разными предприятиями чтобы разместить датчики, собрать данные в разных помещениях с разными материалами.

Едва ли не самым важным фактором оказывается качество, разнообразие и природа

данных в распоряжении исследователей. Большой объём информации – необходимое, но отнюдь не достаточное условие для применения ИИ. Не любой большой массив данных можно обрабатывать с помощью обучающихся алгоритмов, с этим столкнулись все команды, с которыми мы работали. Характер данных, их машиночитаемость, структурированность и однородность напрямую определяют возможность применения ИИ. Один из представителей центров ИИ говорит, что «гигиена данных» – ядро всей работы. Неподготовленные данные требуют предварительной, зачастую нетривиальной обработки, что может заблокировать проект или сделать ИИ не инструментом анализа, а инструментом очистки данных. Кейсы историков (низкое качество распознавания текста) и теплофизиков (несопоставимость *Excel*-таблиц) наглядно демонстрируют это. Успех проекта теплофизиков стал возможен благодаря важному решению, найденному ИИ-специалистами, которым удалось убедить учёных отказаться от привычных для них форм данных в пользу необработанного источника (базы контроллера). Это переворот в культуре проведения исследования, который задаёт множество эпистемических и политических вопросов: кто теперь определяет качество материалов, кто отвечает за методы и за результаты научной работы. У физиков теряет ценность эксперимент и анализ его результатов, которые нужно привести в формат научного текста. У историков также меняется отношение в работе с первоисточником, которая опосредуется техническими средствами. Далеко не все научные коллективы готовы принять новую эпистему. Гигиена данных при отчуждении работы учёного от них – ядро всей работы.

Чтобы обеспечить чистоту данных, исследователи перестраивают работу либо сами, либо в результате принуждения. Это позволяет работать с большими массивами данных программными средствами.

Третий фактор имеет организационную природу. Междисциплинарное сотрудниче-

ство между учёными-предметниками и специалистами в области компьютерных наук почти всегда необходимо для внедрения ИИ. Это не просто обмен услугами, а формирование культуры, где различные типы экспертизы должны работать на уровне повседневных практик. Критически важна готовность и возможность выстроить взаимодействие.

Некоторые, не связанные с техническими ограничениями, барьеры могут блокировать возможность применения ИИ, не позволяя сформировать необходимый массив данных. Они актуализируют институциональные и этические вопросы. Ярче всего это проявляется в медико-биологических науках, где ограниченный доступ к объекту исследования и строгие этико-бюрократические фильтры создают практически непреодолимые барьеры.

Наличие институциональной инфраструктуры и доступ к компетенциям

Чтобы внедрить ИИ, нужно перераспределить то, что К. Кнорр-Цетина называет «не-человеческой социальностью» [3], связи между учёными, инструментами, данными и институтами. Ни один из коллективов не обладал изначально компетенциями в разработке и использовании передовых цифровых технологий, но все они имели некоторый опыт построения такой социальности для разных проектов. Это позволило в разной степени для каждого случая выстроить равновесие имеющихся и недостающих компетенций.

При планировании исследования с применением ИИ коллективы обращают внимание на выполнение нескольких условий. Во-первых, доступ к специалистам, которые обладают непривычными для исследователей компетенциями. Во-вторых, возможность платить специалистам и долгосрочной перспективе. Центры ИИ должны решить эту задачу. Они созданы в рамках крупных организаций: НИИ и университетов. Все четыре лаборатории взаимодействовали с такими центрами. Один из сотрудников центра ИИ описывал модель его работы сле-

дующим образом: «...то, что делают большие компании – для широкой аудитории. Всегда найдутся люди, у которых запросы более узкие. Продукт, который сделан под широкую аудиторию, может не подходить им, и слишком дорого просить компании сузить рамки работы их продукта. Вот они и идут к нам». Стратегия Центра нацелена на реализацию небольших и уникальных проектов, которые характерны для научной среды. Информанты из числа датасайентистов отмечали личный интерес в работе над такими проектами, с точки зрения как профессионального роста, так и участия в научной деятельности. Это позволяет несколько снизить стоимость работы ИИ-специалистов.

Недостаточно только создать центры ИИ. Работа ИИ-специалистов зачастую вызывают скепсис из-за сложности валидации результатов научных исследований. Учёные видят себя носителями дисциплинарных методов и правил, требующих привычных для их науки подтверждений. Вмешательство, нарушающее физические или интерпретативные принципы, воспринимается как угроза эпистемической автономии. Существует общепринятая практика среди медиков и физиков экспериментально подтверждать результаты математических расчётов и прогнозов ИИ: «Матмодель, реализованная в компьютере, может насчитать всё что угодно. Верификация – это всегда эксперимент. ...Как бы там ни пытались строить виртуальные лаборатории, ты новое в них не сделаешь» (фотонщик 2). Об этом же говорят теплофизики. ИИ серьёзно расширяет возможности для моделирования, но сохраняется необходимость доказать связь расчётов с эмпирическими данными. Подтверждающий эксперимент создаёт уверенность в том, что можно брать ответственность за расчёты ИИ. С другой стороны, коллектив историков не имел возможности экспериментально проверить вычисления: «...Почему вы выбрали именно эти слова, а не эти? Как это работает? Какие алгоритмы? Как это всё возможно? Я могу миллион во-

просов перечислить. Большинство моих коллег именно так и скажут». Здесь мы видим различие эпистемических культур: в естественных науках существует устоявшаяся практика экспериментальной верификации, тогда как в исторической культуре, ориентированной на интерпретацию, вопрос о верификации алгоритмических результатов оказывается принципиально открытым, что создаёт недоверие к методу.

Руководители крупных организаций проявляют инициативу, принуждают скептиков к инновациям, устраивая совместные мероприятия для знакомства с новым инструментом, а в некоторых организациях даже внедряют внутреннее финансирование таких проектов. На таких встречах в двух случаях было принято решение привлечь ИИ.

Представители центра ИИ пытались заинтересовать исследователей. Даже у теплофизиков, руководитель которых был знаком с руководителем команды ИИ-специалистов, идея применения ИИ возникла в процессе реализации проекта с подачи датасайентистов. Наиболее значимым оказалось то, что ИИ-специалисты смогли предложить стратегию исследования, очевидно связанную с результатом. Таким образом, научные коллективы смогли познакомиться с ИИ во многом благодаря административному давлению.

Иногда инициативу при посредничестве институтов развития, проявляют одновременно и исследователи, и ИИ-специалисты. Руководитель одного из центров ИИ описывал, что не только организации, но и государство проводит хакатоны для внедрения ИИ в научную работу. На соревнованиях кейсы предлагают научные организации, выигравшие решения получают успешную коллаборацию ИИ-специалистов с ведущими учреждениями страны, не только своей организации, а также финансирование научных проектов.

Эффективная коммуникация и «перевод» между дисциплинами

Успешная коллаборация невозможна без общего научного языка. К.Кнорр-Цетина

называет это формированием эпистемической машины, где различные типы знания (предметное и вычислительное) должны сосуществовать на уровне повседневных практик, а не только на уровне формальных договорённостей [3]. Опыт теплофизиков показывает, что этот этап может занимать до полугода, но без него работа ИИ как стратегии исследования невозможно. Необходимо выявить места пересечения экспертиз и научиться понимать терминологию друг друга. Обратное движение – знакомство предметников с ИИ-технологиями – на первом этапе деятельности ИИ практически не происходит. Такая асимметрия соответствует наблюдениям Ш.Трэвик о том, что в науке существует устойчивое распределение ролей между теми, кто работает с приборами/инструментами, и теми, кто задаёт содержательные вопросы; для интеграции ИИ необходим пересмотр разделения [4]. Если руководитель коллектива принимает решение сделать ИИ одним из постоянно применяемых методов (теплофизики и, вероятно, фотонщики), то выстраивается инфраструктура взаимопроникновения методов ИИ в работу коллектива и подготовку кадров в структуре лаборатории. В случае с историками и медицинскими биофизиками не произошло трансформации научной инфраструктуры. Технология осталась на аутсорсе, и с окончанием проекта закончилось применение ИИ. Далее мы подробно рассмотрим причины этого.

Барьеры для внедрения ИИ

Искусственный интеллект может обеспечить преимущества исследованию, однако существует и риски.

Низкое качество и неготовность данных – один из самых серьёзных барьеров. Историки столкнулись с несовершенством распознавания текстов, теплофизики – с разрозненностью и нестандартизированностью таблиц, биофизики – с малым числом оцифрованных медицинских данных. В первых двух случаях это потребовало не тривиальных усилий на предварительную

очистку и подготовку данных, в последнем стало одной из причин, заблокировавших проект.

Институциональные ограничения, в том числе связанные с проблемой верификации результатов, полученных с помощью ИИ, и недоверие к «чёрному ящику» технологии – препятствие для признания метода. В естественных науках (физики, медики) существует устоявшаяся практика экспериментальной проверки: *«Любая математика... может насчитать все, что угодно. Верификация – это всегда эксперимент»*. В гуманитарных науках такой возможности нет. Вопросы коллег: *«Почему вы выбрали именно эти слова, а не эти? Как это работает?»*, – остаются без ответа, что порождает недоверие и формирует осторожное отношение к новому методу. С точки зрения руководителя, ИИ – вспомогательный сервис, но именно историк производит новое знание. По его мнению, сопротивление оправдано тем, что *«ИИ может насчитать что угодно»*, в его дисциплине проверить мы это никак не можем.

Другой историк описывал сложности со стороны журналов: *«...с публикациями возникают сложности. Мы отправляем статью, нам приходит ответ, что у них нет рецензентов по профилю. Или бывают ответы в духе... [статья] вообще им не по профилю»*. Статья оказалась слишком необычной методически, журналы не соглашались принимать её в печать.

Критическим ограничением для большинства проектов оказалось финансирование. У проектов короткие горизонты планирования: *«...грант РНФ, который составляет от полутора миллионов до шести миллионов в год. Этих денег хватает только на то, чтобы какую-то одну молекулу испытать»* (биофизик 1). Государственное финансирование позволяет исследовать только небольшой объём эмпирического материала, поэтому скорость набора данных, на основе которых можно обучать ИИ, невысока, возможности привлекать сторонних экспертов ограничены.

Обсуждение

В обзоре литературы мы упоминали, что корпус литературы, посвящённой исследуемой теме, ограничен. Интеграция искусственного интеллекта как инфраструктуры в научные исследования была исследована Д. Лю и Х.В. Джагадишем [34], С. Бьянкини и коллегами [35]. Результаты исследования демонстрируют как соответствие, так и расхождения с позициями ключевых авторов, работающих над проблемой включения ИИ в научные исследования.

Д. Лю и Х.В. Джагадиш [34] описывают материальную составляющую ИИ: вычислительные центры и обучение исследователей как необходимые условия для внедрения технологии. Наше исследование подтверждает значимость этого фактора, но с существенными уточнениями. Коллектив биофизиков действительно столкнулся с отсутствием локальной вычислительной инфраструктуры, что стало одним из препятствий к применению нейросетей в их исследовании. Учёные могут компенсировать технические ограничения распределёнными сетями сотрудничества, благодаря которым можно преодолеть географическую разобщённость. Например, Центр ИИ стал одним из победителей в хакатоне, что позволило ему помочь организации из другого региона решить научную задачу. Команда теплофизиков тоже начала участвовать в нетипичных научных проектах. Сетевое взаимодействие может быть альтернативой локальной инфраструктуре. Однако такой подход работает, когда есть временные, финансовые и организационные ресурсы для поддержания этих связей – фактор, который не выделяют ключевые публикации [34; 35].

Предположение Д. Лю и Х. В. Джагадиша о необходимости обучать исследователей навыкам ИИ [34] требует переформулировки. В исследованных лабораториях успешная интеграция ИИ произошла не через прямое обучение учёных работе с нейросетями, а через выстраивание устойчивой коммуникации между двумя экспертными сообще-

ствами. В кейсе теплофизиков ключевую роль сыграл руководитель лаборатории как «переводчик», способный формулировать научные задачи на языке, понятном программистам. Это не обучение, а постепенная адаптация специалистов по ИИ к логике дисциплины. Компетенции формируются как результат сотрудничества, при котором обе стороны обогащают экспертизы друг друга. Руководитель теплофизиков включил обучение программированию и работе с нейросетями для начинающих исследователей уже после успешной коллаборации.

Возможность учёных перенять новые исследовательские практики во многом связана с локальной традицией подготовки научных специалистов. Теплофизики и биофизики, ранее работавшие с более простыми инструментами (*Excel* и *SPSS*), столкнулись с трудностями при интеграции нового метода. В первом случае база данных, обработанных в *Excel*, оказалась несовместима с ИИ. В случае с биофизиком, работавшим с радиомикой, многие возможности технологии оказались недоступны из-за отсутствия возможности достаточно тщательной переподготовки.

С. Бьянкини и коллеги [35] подчёркивают влияние коллективов, использующих ИИ и доли молодых учёных на внедрение технологии. Наши данные показывают, что ключевой фактор – инициатива руководства, которая помогает преодолеть скепсис учёных. Мы не наблюдали вариаций организационных культур коллективов. Во всех случаях присутствуют учёные-предметники и программисты, выполнявшие соответствующие задачи. В научных группах, вступающих в длительные коллаборации, специалисты по ИИ могут полностью влиться в коллектив. В краткосрочных проектах, программисты выступают, скорее, как внешние подрядчики. Коллективы, столкнувшиеся с трудностями при внедрении ИИ, отмечали большее количество институциональных преград: этические комитеты в случае с биофизиками или методический консерватизм научного сообщества историков.

Д. Лю и Х.В. Джагадиш, рассуждая в основном о генеративном ИИ в академической среде [34], не связывают исследовательскую задачу и выбор методологии. Наши данные показывают, что гармонизация характера задачи и формата данных необходимы для внедрения ИИ. Технология вводится, когда исследовательская задача становится технически нерешаемой другими средствами или когда позволяет сформулировать проблему, которая ранее не могла быть решена. Ещё одним стимулом для внедрения новых технологий становится недостаточное технологическое оснащение в условиях санкций и ограничений международных контактов, и тогда учёные применяют новые технологии для покрытия складывающегося дефицита ресурсов [41]. Неудача гармонизации задачи и данных – критический барьер для интеграции ИИ в исследование. Может возникать противоречие между результатами работы ИИ и дисциплинарными представлениями о качестве научного знания. Эту проблему описывали два исследовательских коллектива. Группа историков столкнулась с недоверием научного сообщества к результатам, полученным с помощью ИИ. Качество научного знания в этом случае ассоциировалось с методом, которым оно было получено. Учёные обошли проблему, публикуя результаты исследований в междисциплинарных журналах и изданиях смежных дисциплин. С другой стороны, биофизикам не удалось разрешить это противоречие. В отличие от историков, в их случае сопротивление шло «изнутри» исследовательской группы и происходило на уровне техники сбора и анализа данных, а не метода. Информант описывал фрустрацию из-за невозможности понять принципы, по которым ИИ производит вычисления, отсутствия очевидной связи с эмпирической составляющей биомедицины. Соответствие получаемого знания дисциплинарным нормам – важный фактор принятия технологии, но не решающий.

Таким образом, существуют три взаимосвязанных условия, выявленных при эм-

пирическом анализе кейсов научных лабораторий. Гармонизация исследовательской задачи и формата данных предполагает, что ИИ интегрируется тогда, когда традиционные методы становятся технически неэффективными, либо появляется возможность решать новые проблемы. Активная организационная политика позволяет преодолеть начальный скепсис и сформировать необходимые структуры. Успешная коммуникация между исследователями и специалистами по ИИ проходит через «переводчиков» (например, руководителей лабораторий), обеспечивая взаимное обогащение компетенций без прямого обучения. Гармонизация задач и данных предполагает, что искусственный интеллект возможно обучить на имеющемся массиве. Активная организационная политика должна преодолеть дисциплинарный консерватизм. Наконец, для возможности коммуникации между специалистами необходимо обеспечить финансирование проекта. Реализация факторов и преодоление барьеров должны происходить последовательно и итеративно. Сформулированная цель позволяет заняться организацией коллектива, а формирование общего видения открывает возможности для коммуникации с ИИ-специалистами, что может снова привести к необходимости уточнять задачу.

Заключение

В исследовании были выявлены ключевые факторы, определяющие интеграцию искусственного интеллекта в деятельность научных коллективов. На основе анализа четырёх кейсов российских лабораторий показано, что внедрение ИИ не столько технологически обусловленный процесс, сколько многоуровневое согласование исследовательской задачи, имеющихся ресурсов, компетенций, организационной культуры и дисциплинарных норм валидации полученного знания. Неудача включения ИИ в исследование часто связана с разрывом хотя бы в одном звене этой цепи.

Концептуализация эмпирических данных показала, что эпистемическая культура фильтрует практики, определяя их как развитие существующих или как чужеродное влияние. Этот фильтр не детерминирует исход. Агентство ИИ опосредована конфигурацией практик научной работы коллектива, а не только технологической, дисциплинарной или институциональной рамкой. Именно это определяет, станет ли ИИ внешним ресурсом или частью трансформированной эпистемической модели. Устойчивая интеграция происходит только тогда, когда все три измерения работают вместе. Эпистемическая культура не блокирует метод, а коллектив понимает, как гармонизировать исследовательскую задачу и форматы данных. Сложившиеся практики работы научной группы обеспечивают долгосрочную материальную, временную и компетентностную поддержки. Организационная политика помогает достичь уровня, необходимого для преодоления границ между ИИ-отраслью и научными дисциплинами.

Исследование обладает рядом преимуществ. Мы проводим триангуляцию через множественное кейс-стади. Такой подход обеспечивает надёжность выводов путём привлечения представителей разных дисциплин и разных степеней успеха внедрения ИИ. Фокус на факторы, влияющие на принятие решений при отслеживании реального опыта коллективов, позволяет выйти за пределы «солюционистских» нарративов, доминирующих в литературе об ИИ.

Вместе с тем исследование основано на небольшой выборке, что ограничивает возможность обобщения результатов на все российские научные коллективы. Эмпирический материал на основе выборки представителей разных дисциплин позволяет делать широкие обобщения, но в процессе теряется значительная часть дисциплинарной специфики. Работа открывает перспективы для будущих исследований. Расширение географии и масштаба позволит выявить, степень универсальности выводов. Лонгитюдное ис-

следование позволит проследить развитие коллективов, которые интегрировали ИИ, чтобы понять, как меняется характер их работы. Анализ на уровне отдельных организаций позволит исследовать, какие управленческие стратегии наиболее эффективны.

Литература

1. Попова Е.В. Нейросеть для науки в действии: невидимая работа и пересборка практик теплофизической лаборатории // Этнографическое обозрение. – 2026. – № 2.
2. Селянин Я.В. Государственная политика США в области искусственного интеллекта: цели, задачи, перспективы реализации // Проблемы национальной стратегии. – 2020. – № 4. – С. 140–163. – EDN: FOAPWY.
3. Knorr-Cetina K.D. The Scientist as a Socially Situated Reasoner: From Scientific Communities to Transscientific Fields / The manufacture of knowledge: An essay on the constructivist and contextual nature of science. – Elsevier, 2013. – P. 68–93. – ISBN: 9781483285740.
4. Traweek S. Beamtimes and lifetimes: The world of high energy physicists. – Harvard University Press, 1992. – ISBN: 978-0674063488.
5. Neumann M., Rauschenberger M., Schüp E. M. “We need to talk about ChatGPT”: The future of AI and higher education // 2023 IEEE/ACM 5th International Workshop on Software Engineering Education for the Next Generation (SEENG). IEEE, 2023. – P. 29–32. – DOI: 10.1109/SEENG59157.2023.00010.
6. Кошкина Е.А., Бордовская Н.В., Гнедых Д.С., Хромова М.А., Демьянчук Р.В. и др. Генеративный искусственный интеллект в высшем образовании: обзор теоретических подходов и практик применения // Высшее образование в России. – 2025. – Т. 34, № 6. – С. 36–57. – DOI: 10.31992/0869-3617-2025-34-6-36-57.
7. Bond M., Khosravi H., de Laat M., Bergdahl N. et al. A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigor // International journal of educational technology in higher education. – 2024. – Vol. 21, no. 4. – 41 p. – DOI: 10.1186/s41239-023-00436-z.
8. Kim J., Klopfer M., Grohs J.R., Eldardiry H., Weichert J. et al. Examining faculty and student perceptions of generative AI in university courses // Innovative Higher Education. – 2025. – Vol. 50. – P. 1–33. – DOI: 10.1007/s10755-024-09774-w.
9. Xu Y., Liu X., Cao X., Huang Ch., Liu E. et al. Artificial intelligence: A powerful paradigm for scientific research // The Innovation. – 2021. – Vol. 2, no. 4. – 21 p. – DOI: 10.1016/j.xinn.2021.100179.
10. Grossmann I., Feinberg M., Parker D.C., Christakis N.A., Tetlock P.E., Cunningham W.A. AI and the transformation of social science research // Science. – 2023. – Vol. 380, no. 6650. – P. 1108–1109. – DOI: 10.1126/science.adi1778.
11. Van Noorden R., Perkel J.M. AI and science: what 1,600 researchers think // Nature. – 2023. – Vol. 621. – P. 672–675. – DOI: 10.1038/d41586-023-02980-0.
12. Berens P., Cranmer K., Lawrence N.D., von Luxburg U., Montgomery J. AI for science: an emerging agenda // Dagstuhl Seminar 22382 “Machine Learning for Science: Bridging Data-Driven and Mechanistic Modelling”. – 2023. – 44 p. – arXiv preprint arXiv:2303.04217.
13. Cheng K., Li Zh., He Y., Guo Q., Lu Y. et al. Potential use of artificial intelligence in infectious disease: take ChatGPT as an example // Annals of biomedical engineering. – 2023. – Vol. 51, no. 6. – P. 1130–1135. – DOI: 10.1007/s10439-023-03203-3.
14. Reddy C.K., Shojae P. Towards scientific discovery with generative AI: Progress, opportunities, and challenges // Proceedings of the AAAI Conference on Artificial Intelligence. – 2025. – Vol. 39, no. 27. – P. 28601–28609. – DOI: 10.1609/aaai.v39i27.35084.
15. Latif E., Parasuraman R., Zhai X. Physics assistant: An LLM-powered interactive learning robot for physics lab investigations // 2024 33rd IEEE International Conference on Robot and Human Interactive Communication (ROMAN). IEEE, 2024. – P. 864–871. – DOI: 10.1109/ROMAN60168.2024.10731312.
16. Gao J., Choo K., Cao J., Lee R.K.-W. et al. CoAI-coder: Examining the effectiveness of AI-assisted human-to-human collaboration in qualitative analysis // ACM Transactions on Computer-Human Interaction. – 2023. – Vol. 31, no. 1. – P. 1–38. – DOI: 10.1145/3617362.
17. Sarker I.H. AI-based modeling: techniques, applications and research issues towards automation, intelligent and smart systems // SN computer science. – 2022. – Vol. 3, no. 2. – Article no. 158. – DOI: 10.1007/s42979-022-01043-x.

18. Wagner G., Lukyanenko R., Парй G. Artificial intelligence and the conduct of literature reviews // *Journal of Information Technology*. – 2022. – Vol. 37, no. 2. – P. 209–226. – DOI: 10.1177/02683962211048201.
19. Vaccaro M., Almaatouq A., Malone T. When combinations of humans and AI are useful: A systematic review and meta-analysis // *Nature Human Behaviour*. – 2024. – Vol. 8, no. 12. – P. 2293–2303. – DOI: 10.1038/s41562-024-02024-1.
20. Телицына А.Ю. Оптимизация научной деятельности через интеграцию ИИ: нейронные сети как инструмент в работе с академической литературой // *Мониторинг*. – 2024. – Т. 183, № 5. – С. 218–236. – DOI: 10.14515/monitoring.2024.5.2623.
21. Oduoye M.O., Javed B., Gupta N., Sih C.M.V. et al. Algorithmic bias and research integrity; the role of nonhuman authors in shaping scientific knowledge with respect to artificial intelligence: a perspective // *International Journal of Surgery*. – 2023. – Vol. 109, no. 10. – P. 2987–2990. – DOI: 10.1097/JS9.0000000000000552.
22. Brodie M.L. Future intelligent information systems: AI and database technologies working together // *Readings in artificial intelligence and database*. Morgan Kaufman, 1989. – P. 623–641. – DOI: 10.1016/B978-0-934613-53-8.50048-0.
23. Bryda G., Sadowski D. From words to themes: AI-powered qualitative data coding and analysis // *World conference on qualitative research*. – Cham: Springer Nature Switzerland, 2024. – P. 309–345. – DOI: 10.1007/978-3-031-65735-1_19.
24. Khalifa M., Albadawy M. Using artificial intelligence in academic writing and research: An essential productivity tool // *Computer Methods and Programs in Biomedicine Update*. – 2024. – Vol. 5. – DOI: 10.1016/j.cmpbup.2024.100145.
25. Longkumer T. Generative AI in Anthropology: Redefining Fieldwork, Textualisation, and Collaborative Knowledge Production // *Teaching Anthropology*. – 2025. – Vol. 14, no. 2. – P. 118–130. – DOI: 10.22582/ta.v14i2.778.
26. Chakravorti T., Wang X., Venkit P.N., Koneru S., Munger K. et al. Social Scientists on the Role of AI in Research // *arXiv preprint*. – 2025. – 14 p. – DOI: 10.48550/arXiv.2506.11255.
27. Sarker I.H. AI-based modeling: techniques, applications and research issues towards automation, intelligent and smart systems // *SN computer science*. – 2022. – Vol. 158, no. 3. – 20 p. – DOI: 10.1007/s42979-022-01043-x.
28. Nagar S. Artificial Intelligence in Scientific Research: Lessons for SPIs // *Science-Policy Brief for the Multistakeholder Forum on Science, Technology and Innovation for the SDGs*, May 2024. United Nations. 2024. – 4 p. – URL: <https://sdgs.un.org/documents/nagar-s-artificial-intelligence-scientific-research-lessons-spis-55347> (дата обращения: 09.12.25).
29. Mohammadi E., Cai Y., Novin A., Vera V., Soltanmohammadi E. Who is a scientist? Gender and racial biases in google vision AI /// *AI and Ethics*. – 2025. – Vol. 5. – P. 4993–5010. – DOI: 10.1007/s43681-025-00742-4.
30. Messeri L., Crockett M.J. Artificial intelligence and illusions of understanding in scientific research // *Nature*. – 2024. – Vol. 627. – P. 49–58. – DOI: 10.1038/s41586-024-07146-0.
31. Kobak D., González-Márquez R., Horvát E.-Á., Lause J. Delving into LLM-assisted writing in biomedical publications through excess vocabulary // *Science Advances*. – 2025. – Vol. 11, no. 27. – 8 p. – DOI: 10.1126/sciadv.adt3813.
32. Watermeyer R., Lanclos D., Phipps L., Shapiro H., Guizzo D. et al. Academics' Weak(ening) Resistance to Generative AI: The Cause and Cost of Prestige? // *Postdigital Science and Education*. – 2025. – Vol. 7. – P. 1171–1191. – DOI: 10.1007/s42438-024-00524-x.
33. Carobene A., Padoan A., Cabitza F., Banfi G., Plebani M. Rising adoption of artificial intelligence in scientific publishing: evaluating the role, risks, and ethical implications in paper drafting and review process // *Clinical Chemistry and Laboratory Medicine (CCLM)*. – 2024. – Vol. 62, no. 5. – P. 835–843. – DOI: 10.1515/cclm-2023-1136.
34. Liu J., Jagadish H.V. Institutional efforts to help academic researchers implement generative AI in research // *Harvard Data Science Review*. – 2024. – No. 5. – 25 p. – DOI: 10.1162/99608f92.2c8e7e81.
35. Bianchini S., Möller M., Pelletier P. Drivers and barriers of ai adoption and use in scientific research // *Technological Forecasting and Social Change*. – 2025. – Vol. 220. – 16 p. – DOI: 10.1016/j.techfore.2025.124303.
36. Prieto-Gutierrez J. J., Segado-Boj F., Franzá F.D.S. Artificial intelligence in social science: A study based on bibliometrics analysis // *Human Technology*. – 2023. – Vol. 19, no. 2. – P. 149–162. – DOI: 10.48550/arXiv.2312.10077.
37. Попова Е.В., Мацепуро Д.М. Исследовательская этика: представления и практики российской

- ских молодых учёных // Высшее образование в России. – 2024. – Т. 33, № 7. – С. 124–143. – DOI: 10.31992/0869-3617-2024-33-7-124-143.
38. Steinhart I., Bauer M., Wünsche H., Schimmmler S. The connection of open science practices and the methodological approach of researchers // *Quality & Quantity*. – 2023. – Vol. 57. – no. 4. – P. 3621–3636. – DOI: 10.1007/s11135-022-01524-4.
 39. Bennett D.S., Stam A.C. Research design and estimator choices in the analysis of interstate dyads: When decisions matter // *Journal of conflict resolution*. – 2000. – Vol. 44, no. 5. – P. 653–685. – DOI: 10.1177/0022002700044005005.
 40. Парамзин Ф.Н., Какоткин В.В., Буркин Д.А., Агапов М.А. Радиомика и искусственный интеллект в дифференциальной диагностике опухолевых и неопухолевых заболеваний поджелудочной железы (обзор) // Хирургическая практика. – 2023. – № 1. – С. 53–65. – DOI: 10.38181/2223-2427-2023-1-5.
 41. Попова Е.В. Исследовательские инфраструктуры российских лабораторий: до и после санкций // *Этнографическое обозрение*. – 2025. – № 6. – С. 168–191. – DOI: 10.7868/S3034627425060109.
- Благодарности.** Результаты были получены в рамках выполнения государственного задания Минобрнауки России, проект № FSWM-2024-0008ю.
- Статья поступила в редакцию 03.02.2026*
Принята к публикации 10.04.2026

References

1. Popova, E.V. (2026). Neuroset' dlia nauki v deistvii: nevidimaia rabota i peresborka praktik teplofizicheskoi laboratorii [Neural network for Science in action: invisible work and reassembly of thermal Physics Laboratory practices]. *Etnograficheskoe obozrenie* [Ethnographic Review]. No. 2. (In print).
2. Selyanin, Ya.V. (2020). The U.S. National Policy In The Area Of Artificial Intelligence: Aims, Objectives And Prospects For Realization. *Problemy natsional'noi strategii = National Strategy Issues*. No. 4, pp. 140–163. Available at: https://elibrary.ru/download/elibrary_44324190_37383096.pdf (accessed 01.11.2025). (In Russ., abstract in Eng.).
3. Knorr-Cetina, K.D. (2013). The Scientist as a Socially Situated Reasoner: From Scientific Communities to Transscientific Fields. In: *The Manufacture of Knowledge: An Essay on the Constructivist and Contextual Nature of Science*. Elsevier. P. 68–93. ISBN: 9781483285740.
4. Traweek, S. (1992). *Beamtimes and Lifetimes: The World of High Energy Physicists*. Harvard University Press. ISBN: 978-0674063488.
5. Neumann, M., Rauschenberger, M., Schün, E.M. (2023). “We Need to Talk about ChatGPT”: The Future of AI and Higher Education. In: *2023 IEEE/ACM 5th International Workshop on Software Engineering Education for the Next Generation (SEENG)*. IEEE, pp. 29–32, doi: 10.1109/SEENG59157.2023.00010.
6. Koshkina E.A., Brodovskaya, N.V., Gnedykh, D.S., Khromova, M.A., Demyanchuk, R.V. et al. (2025) Generative Artificial Intelligence in Higher Education: A Review of Theoretical Approaches and Application Practices. *Vysshee obrazovanie v Rossii = Higher Education in Russia*. Vol. 34, no. 6, pp. 36–57, doi: 10.31992/0869-3617-2025-34-6-36-57 (In Russ. Abstract in Eng.).
7. Bond, M., Khosravi, H., de Laat, M., Bergdahl, N. et al. (2024) A Meta Systematic Review of Artificial Intelligence in Higher Education: A Call for Increased Ethics, Collaboration, and Rigor. *International Journal of Educational Technology in Higher Education*. Vol. 21, no. 4, 41 p., doi: 10.1186/s41239-023-00436-z.
8. Kim J., Klopfer, M., Grohs, J.R., Eldardiry, H., Weichert, J. et al. (2025) Examining Faculty and Student Perceptions of Generative AI in University Courses. *Innovative Higher Education*. Vol. 50, pp. 1–33, doi: 10.1007/s10755-024-09774-w.
9. Xu Y., Liu, X., Cao, X., Huang, Ch., Liu, E. et al. (2021). Artificial Intelligence: A Powerful Paradigm for Scientific Research. *The Innovation*. Vol. 2, no. 4, 21 p., doi: 10.1016/j.xinn.2021.100179.
10. Grossmann, I., Feinberg, M., Parker, D.C., Christakis, N.A., Tetlock, P.E., Cunningham, W.A. (2023). AI and the Transformation of Social Science Research. *Science*. Vol. 380, no. 6650, pp. 1108–1109, doi: 10.1126/science.adi1778.

11. Van Noorden, R., Perkel, J.M. (2023) AI and Science: What 1,600 Researchers Think. *Nature*. Vol. 621, pp. 672-675, doi: 10.1038/d41586-023-02980-0.
12. Berens, P., Cranmer, K., Lawrence, N.D., von Luxburg, U., Montgomery, J. (2023). AI for Science: An Emerging Agenda. In: *Dagstuhl Seminar 22382 "Machine Learning for Science: Bridging Data-Driven and Mechanistic Modelling"*. 44 p. arXiv preprint arXiv:2303.04217.
13. Cheng, K., Li, Zh., He, Y., Guo, Q., Lu, Y. et al. (2023). Potential Use of Artificial Intelligence in Infectious Disease: Take ChatGPT as an Example. *Annals of Biomedical Engineering*. Vol. 51, no. 6, pp. 1130-1135, doi: 10.1007/s10439-023-03203-3.
14. Reddy, C.K., Shojaee, P. (2025). Towards Scientific Discovery with Generative AI: Progress, Opportunities, and Challenges. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. Vol. 39, no. 27, pp. 28601-28609, doi: 10.1609/aaai.v39i27.35084.
15. Latif, E., Parasuraman, R., Zhai, X. (2024). Physics Assistant: An LLM-Powered Interactive Learning Robot for Physics Lab Investigations. In: *2024 33rd IEEE International Conference on Robot and Human Interactive Communication (ROMAN)*. IEEE. Pp. 864-871, doi: 10.1109/RO-MAN60168.2024.10731312.
16. Gao, J., Choo, K., Cao, J., Lee, R.K.-W. et al. (2023) CoAICoder: Examining the Effectiveness of AI-Assisted Human-to-Human Collaboration in Qualitative Analysis. *ACM Transactions on Computer-Human Interaction*. Vol. 31, no. 1, pp. 1-38, doi: 10.1145/3617362.
17. Sarker, I.H. (2022). AI-based Modeling: Techniques, Applications and Research Issues Towards Automation, Intelligent and Smart Systems. *SN Computer Science*. Vol. 3, no. 2, article no. 158, doi: 10.1007/s42979-022-01043-x.
18. Wagner, G., Lukyanenko, R., Paré, G. (2023) Artificial Intelligence and the Conduct of Literature Reviews. *Journal of Information Technology*. Vol. 37, no. 2, pp. 209-226, doi: 10.1177/02683962211048201.
19. Vaccaro, M., Almaatouq, A., Malone, T. (2024). When Combinations of Humans and AI Are Useful: A Systematic Review and Meta-Analysis. *Nature Human Behaviour*. Vol. 8, no. 12, pp. 2293-2303, doi: 10.1038/s41562-024-02024-1.
20. Telitsyna, A.Yu. (2024) Optimization of Scientific Activities Through AI Integration: Neural Networks as a Tool in Working with Academic Literature. *Monitoring = Monitoring*. Vol. 183, no. 5, pp. 218-236, doi: 10.14515/monitoring.2024.5.2623. (In Russ., abstract in Eng.).
21. Oduoye, M.O., Javed, B., Gupta, N., Sih, C.M.V. (2023). Algorithmic Bias and Research Integrity: The Role of Nonhuman Authors in Shaping Scientific Knowledge with Respect to Artificial Intelligence: A Perspective. *International Journal of Surgery*. Vol. 109, no. 10, pp. 2987-2990, doi: 10.1097/JS9.0000000000000552..
22. Brodie, M.L. (1989). Future Intelligent Information Systems: AI and Database Technologies Working Together. *Readings in Artificial Intelligence and Databases*. Morgan Kaufman. Pp. 623-641, doi: 10.1016/B978-0-934613-53-8.50048-0.
23. Bryda, G., Sadowski, D. (2024). From Words to Themes: AI-Powered Qualitative Data Coding and Analysis. In: *World Conference on Qualitative Research*. Cham : Springer Nature Switzerland. Pp. 309-345, doi: 10.1007/978-3-031-65735-1_19.
24. Khalifa, M., Albadawy, M. (2024). Using Artificial Intelligence in Academic Writing and Research: An Essential Productivity Tool. *Computer Methods and Programs in Biomedicine Update*. Vol. 5, 11 p., doi: 10.1016/j.cmpbup.2024.100145.
25. Chakravorti, T., Wang, X., Venkit, P.N., Koneru, S., Munger, K. et al. (2025) Social Scientists on the Role of AI in Research. *arXiv preprint*. 14 p., doi: 10.48550/arXiv.2506.11255.
26. Sarker, I.H. (2022). AI-based Modeling: Techniques, Applications and Research Issues Towards Automation, Intelligent and Smart Systems. *SN Computer Science*. Vol. 158, no. 3, 20 p., doi: 10.1007/s42979-022-01043-x.
27. Longkumer, T. (2025). Generative AI in Anthropology: Redefining Fieldwork, Textualisation, and Collaborative Knowledge Production. *Teaching Anthropology*. Vol. 14, no. 2, pp. 118-130, doi: 10.22582/ta.v14i2.778.

28. Nagar, S. (2024). Artificial Intelligence in Scientific Research: Lessons for SPIs. In: *Science-Policy Brief for the Multistakeholder Forum on Science, Technology and Innovation for the SDGs, May 2024. United Nations*. 4 p. Available at: <https://sdgs.un.org/documents/nagar-s-artificial-intelligence-scientific-research-lessons-spis-55347> (accessed 01.11.2025).
29. Mohammadi, E., Cai, Y., Novin, A., Vera, V., Soltanmohammadi, E. (2025). Who Is a Scientist? Gender and Racial Biases in Google Vision AI. *AI and Ethics*. Vol. 5, pp. 4993-5010, doi: 10.1007/s43681-025-00742-4.
30. Messeri, L., Crockett, M. J. (2024). Artificial Intelligence and Illusions of Understanding in Scientific Research. *Nature*. Vol. 627, pp. 49-58, doi: 10.1038/s41586-024-07146-0.
31. Kobak, D., González-Márquez, R., Horvát, E.-Á., Lause, J. (2025). Delving into LLM-Assisted Writing in Biomedical Publications Through Excess Vocabulary. *Science Advances*. Vol. 11, no. 27, 8 p., doi: 10.1126/sciadv.adt3813.
32. Watermeyer, R., Lanclos, D., Phipps, L., Shapiro, H., Guizzo, D. et al. (2025). Academics' Weak(ening) Resistance to Generative AI: The Cause and Cost of Prestige? *Postdigital Science and Education*. Vol. 7, pp. 1171-1191, doi: 10.1007/s42438-024-00524-x.
33. Carobene, A., Padoan, A., Cabitza, F., Banfi, G., Plebani, M. (2024). Rising Adoption of Artificial Intelligence in Scientific Publishing: Evaluating the Role, Risks, and Ethical Implications in Paper Drafting and Review Process. *Clinical Chemistry and Laboratory Medicine (CCLM)*. Vol. 62, no. 5, pp. 835-843, doi: 10.1515/cclm-2023-1136.
34. Liu, J., Jagadish, H.V. (2024). Institutional Efforts to Help Academic Researchers Implement Generative AI in Research. *Harvard Data Science Review*. No. 5, 25 p., doi: 10.1162/99608f92.2c8e7e81.
35. Bianchini, S., Møller, M., Pelletier, P. (2025). Drivers and Barriers of AI Adoption and Use in Scientific Research. *Technological Forecasting and Social Change*. Vol. 220, 16 p., doi: 10.1016/j.techfore.2025.124303.
36. Prieto-Gutierrez, J.J., Segado-Boj, F., Franza, F.D.S. (2023). Artificial Intelligence in Social Science: A Study Based on Bibliometrics Analysis. *Human Technology*. Vol. 19, no. 2, pp. 149-162, doi: 10.48550/arXiv.2312.10077.
37. Popova, E.V., Matsepuro, D.M. (2024). Research Ethics: Ideas and Practices of Russian Young Scientists. *Vyshee Obrazovanie v Rossii = Higher Education in Russia*. Vol. 33, no. 7, pp. 124-143, doi: 10.31992/0869-3617-2024-33-7-124-143 (In Russ., abstract in Eng.).
38. Steinhardt, I., Bauer, M., Wünsche, H., Schimmler, S. (2023) The Connection of Open Science Practices and the Methodological Approach of Researchers. *Quality & Quantity*. Vol. 57, no. 4, pp. 3621-3636, doi: 10.1007/s11135-022-01524-4.
39. Bennett, D.S., Stam, A.C. (2000). Research Design and Estimator Choices in the Analysis of Interstate Dyads: When Decisions Matter. *Journal of Conflict Resolution*. Vol. 44, no. 5, pp. 653-685, doi: 10.1177/0022002700044005005.
40. Paramzin, F.N., Kakotkin, V.V., Burkin, D.A., Agapov, M.A. (2023). Radiomics and Artificial Intelligence in The Differential Diagnosis of Tumor and Non-Tumor Diseases of The Pancreas. Review. *Khirurgicheskaya praktika = Surgical Practice*. Vol. 8, no. 1, pp. 53-65, doi: 10.38181/2223-2427-2023-1-5.
41. Popova, E.V. (2025). Research Infrastructures of Russian Laboratories: Before and After Sanctions. *Etnograficheskoe obozrenie = Ethnographic Review*. No. 6, pp. 168-191, doi: 10.7868/S3034627425060109.

Acknowledgement. This research was supported by Ministry of Science and Higher Education of the Russian Federation, project No. FSWM-2024-0008.

*The paper was submitted 03.02.2026
Accepted for publication 10.04.2026*



«Высшее образование в России» – ежемесячный общероссийский научно-педагогический журнал, публикующий результаты фундаментальных, поисковых и прикладных проблемно-ориентированных исследований наличного состояния высшей школы и тенденций её развития, выполненных на стыке наук с позиций педагогики, социологии, истории, экономики и менеджмента. В журнале обсуждаются актуальные вопросы теории и практики модернизации отечественного и зарубежного высшего образования. Особое внимание уделяется проблемам подготовки и повышения квалификации научных и научно-педагогических работников высшей школы.

Целевая аудитория издания – сообщество исследователей и практиков высшего и дополнительного профессионального образования (вузовские и академические учёные, профессорско-преподавательский состав высшей школы, администрация вузов, работники органов управления системой высшего образования, соискатели учёной степени, студенчество). Авторы и читатели журнала – специалисты в области философии образования, педагогики высшей школы, социологии образования.

Миссия журнала – поддержание и развитие единого исследовательского пространства в области наук об образовании в географическом (межрегиональность) и эпистемологическом (междисциплинарность) смысле, а также укрепление межвузовского сотрудничества научно-педагогических работников. Задача – выработка общезначимого языка описания и объяснения современной образовательной реальности, который не только позволяет понимать происходящее, но и сплачивает, объединяет научно-педагогическое сообщество на основе ценностей солидарности, сотрудничества, кооперации и сотворчества.

Журнал входит в Перечень научных изданий, рекомендованных ВАК для публикации результатов исследований по следующим научным специальностям:

- 09.00.08 – Философия науки и техники (философские науки),
- 09.00.11 – Социальная философия (философские науки),
- 13.00.01 – Общая педагогика, история педагогики и образования (педагогические науки),
- 13.00.02 – Теория и методика обучения и воспитания (по областям и уровням образования) (педагогические науки),
- 13.00.08 – Теория и методика профессионального образования (педагогические науки),
- 22.00.04 – Социальная структура, социальные институты и процессы (социологические науки),
- 22.00.06 – Социология культуры (социологические науки)

«Высшее образование в России» публикует теоретические (аналитические, полемические, проблемные) статьи, а также результаты эмпирических и практико-ориентированных исследований, материалы конференций и круглых столов, научные рецензии. В своей деятельности журнал опирается на профессиональные объединения в сфере высшего образования (Российский союз ректоров, Ассоциация технических университетов, Ассоциация инженерного образования России, Ассоциация классических университетов России, Международное общество по инженерной педагогике).

