

АКАДЕМИЧЕСКАЯ ЧЕСТНОСТЬ В ЭПОХУ ГЕНЕРАТИВНОГО ИСКУССТВЕННОГО ИНТЕЛЛЕКТА: ОТ ДЕТЕКЦИИ К ИНСТИТУЦИОНАЛЬНОЙ ПОЛИТИКЕ

Конусова П.*¹, Фартунова М.А.²

¹Независимое агентство аккредитации и рейтинга (IAAR), Астана, Казахстан; ²Софийский университет им. Св. Климента Охридского, София, Болгария

АННОТАЦИЯ. Массовое распространение генеративного искусственного интеллекта (ИИ) поставило под вопрос устоявшиеся механизмы обеспечения академической честности в высшем образовании. Традиционный инструментарий, построенный на выявлении текстовых заимствований, оказался малоприменим к текстам, которые формально являются оригинальными, но не написаны студентом. В статье рассматривается смена содержания понятия академической честности в условиях генеративного ИИ, анализируются ограничения технических средств детекции (высокая доля ложноположительных срабатываний и систематическое смещение против авторов, для которых английский язык не является родным), а также сопоставляются три модели институционального реагирования: запретительная, разрешительно-раскрывающая и интеграционная. Обосновывается тезис о том, что устойчивый ответ на вызов генеративного ИИ лежит не в плоскости детекции, а в плоскости пересмотра дизайна оценивания и закрепления политики использования ИИ в качестве элемента внутренней системы обеспечения качества образования и предмета внешней оценки. Предложен набор практических рекомендаций для вузов Республики Казахстан.

КЛЮЧЕВЫЕ СЛОВА: академическая честность, генеративный искусственный интеллект, обеспечение качества образования, аккредитация, дизайн оценивания, академическое письмо, высшее образование.

Генеративті жасанды интеллект дәуірінде академиялық адалдық: детекциядан институционалдық саясатқа

Конусова П.*¹, Фартунова М.А.²

¹Аккредиттеу және рейтинг тәуелсіз агенттігі (IAAR), Астана, Қазақстан; ²Әулие Климент Охридский атындағы София университеті, София, Болгария.

АНДАТПА. Генеративті жасанды интеллекттің (ЖИ) жаппай таралуы жоғары білім берудегі академиялық адалдықты қамтамасыз етудің қалыптасқан тетіктерін күмән астына алды. Мәтіндік плагиатты анықтауға негізделген дәстүрлі құралдар формальды түрде түпнұсқа болып саналатын, бірақ білім алушының өзі жазбаған мәтіндерге қатысты тиімсіз болып шықты. Мақалада генеративті ЖИ жағдайындағы академиялық адалдық ұғымының мазмұндық өзгеруі қарастырылады, техникалық анықтау құралдарының шектеулері талданады, сондай-ақ институционалдық ден қоюдың үш моделі салыстырылады: тыйым салушы, ашып көрсетуге негізделген рұқсат етуші және интеграциялық. Тұрақты шешім анықтау саласында емес, бағалау дизайнын қайта қарау мен ЖИ-ді пайдалану саясатын білім беру сапасын ішкі қамтамасыз ету жүйесінің элементі әрі сыртқы бағалау нысаны ретінде бекітуде екені негізделеді. Қазақстан жоғары оқу орындарына арналған практикалық ұсынымдар берілген.

ТҮЙІН СӨЗДЕР: академиялық адалдық, генеративті жасанды интеллект, білім беру сапасын қамтамасыз ету, аккредиттеу, бағалау дизайны, академиялық жазу, жоғары білім.

Academic integrity in the age of generative artificial intelligence: from detection to institutional policy

Konussova P.*¹, Fartunova M.A.²

¹Independent Agency for Accreditation and Rating (IAAR), Astana, Kazakhstan; ²Sofia University "St. Kliment Ohridski", Sofia, Bulgaria

ABSTRACT. The rapid diffusion of generative artificial intelligence (AI) has challenged the established mechanisms for safeguarding academic integrity in higher education. Conventional tools built around the detection of textual overlap are of limited use against texts that are formally original yet not written by the student. The article examines how the meaning of academic integrity is shifting under generative AI, analyses the limitations of technical detection tools (high false-positive rates and systematic bias against non-native English writers), and compares three models of institutional response: prohibition, permission with disclosure, and integration. It is argued that a sustainable answer lies not in detection but in redesigning assessment and embedding an institutional AI policy within the internal quality assurance system, thereby making it a subject of external review. A set of practical recommendations for universities of the Republic of Kazakhstan is proposed.

KEYWORDS: academic integrity, generative artificial intelligence, quality assurance, accreditation, assessment design, academic writing, higher education.

ВВЕДЕНИЕ. Появление в открытом доступе больших языковых моделей изменило условия существования высшего образования быстрее, чем система обеспечения качества успела на это отреагировать. Инструмент, способный за несколько секунд сгенерировать связный, стилистически корректный и формально оригинальный текст практически по любой теме, из исследовательской диковинки превратился в повседневное средство, доступное каждому обучающемуся со смартфона. Актуальность темы определяется тем, что этот сдвиг затрагивает не отдельную процедуру проверки, а саму валидность оценивания как способа установить достигнутые результаты обучения.

Проблема носит не столько технический, сколько концептуальный характер. Системы проверки на заимствования, на которых десятилетиями строилась институциональная защита академической честности, отвечают на вопрос «встречался ли этот текст ранее». Генеративная модель производит текст, который ранее не встречался нигде. Формально требование оригинальности соблюдено; фактически работа не отражает ни знаний, ни умений её номинального автора.

Цель статьи — определить, в какой плоскости лежит устойчивый ответ образовательных организаций на вызов генеративного ИИ, и предложить практические рекомендации для организаций высшего образования Республики Казахстан.

Задачи исследования: проследить изменение содержания понятия академической честности в условиях генеративного ИИ; оценить границы применимости технических средств детекции; сопоставить существующие модели институционального реагирования;

обосновать место политики использования ИИ во внутренней системе обеспечения качества образования.

МАТЕРИАЛЫ И МЕТОДЫ ИССЛЕДОВАНИЯ. Объектом исследования выступает институциональное регулирование академической честности в условиях доступности генеративного ИИ, а предметом — модели реагирования образовательных организаций на связанные с ним риски. Исследование носит теоретико-аналитический характер. Использованы следующие методы. Анализ международных нормативных и рекомендательных документов в области академической честности и обеспечения качества образования (рекомендации ЮНЕСКО, позиция Комитета по этике публикаций, Стандарты и рекомендации для обеспечения качества в Европейском пространстве высшего образования) применён для реконструкции формирующейся нормативной рамки. Обзор научных публикаций последних трёх лет использован для оценки эмпирически установленных ограничений средств детекции ИИ-генерируемого текста. Сравнительный анализ институциональных политик университетов, опубликовавших официальные подходы к использованию генеративного ИИ, положен в основу типологии моделей реагирования. Полученные результаты соотнесены с задачами обеспечения качества образования в казахстанском институциональном контексте. Выбор методологического подхода обоснован природой проблемы: поскольку ключевым признаком, отличающим добросовестную работу от недобросовестной, становится не свойство текста, а характер процедуры его создания и раскрытия, адекватным предметом анализа выступают именно институциональные правила, а не показатели технических систем.

РЕЗУЛЬТАТЫ И ИХ ОБСУЖДЕНИЕ. *Смещение содержания понятия академической честности.* Международный центр академической честности определяет академическую честность через набор фундаментальных ценностей — честность, доверие, справедливость, уважение, ответственность и мужество [1]. Существенно, что ни одна из этих ценностей не сформулирована через технологию: они описывают отношение обучающегося к собственной работе, а не способ её создания. Попытка свести проблему генеративного ИИ к вопросу «пользовался ли обучающийся запрещённой программой» изначально смещает фокус. До появления генеративных моделей нарушение академической честности в письменной работе имело, как правило, установившийся след: заимствованный фрагмент существовал где-то ещё и мог быть предъявлен. Работа, выполненная с помощью языковой модели, такого следа не оставляет. Единственный признак, отличающий добросовестную работу от недобросовестной, перемещается из текста в процесс его создания и в характер раскрытия этого процесса.

Отсюда следует первый вывод: в условиях генеративного ИИ академическая честность перестаёт быть свойством текста и становится свойством процедуры. Ключевым оказывается не вопрос «оригинален ли

текст», а вопрос «заявил ли обучающийся, каким образом текст был создан, и соответствует ли это заявление объявленным правилам». Аналогичную позицию заняли и издательские структуры: Комитет по этике публикаций указывает, что инструмент ИИ не может выступать автором, поскольку не способен нести ответственность за содержание, однако факт его использования подлежит раскрытию [2]. Регулирование, таким образом, сместилось от запрета к обязанности раскрытия, и эта логика в равной мере применима к учебному процессу.

Пределы детекции. Первой реакцией многих образовательных организаций стало ожидание технического решения: предполагалось, что вслед за генераторами появятся надёжные детекторы. Накопленные данные эту надежду не подтверждают.

Во-первых, детекторы ИИ-текста работают вероятно и дают значимую долю ложноположительных срабатываний. В условиях, когда решение по академическому проступку затрагивает права обучающегося, вероятностная оценка не может служить достаточным доказательством: обвинение, основанное на статистическом показателе, не является опровержимым для обучающегося, который действительно писал работу самостоятельно. Во-вторых, установлено систематическое смещение таких детекторов против авторов, для которых язык написания не является родным. Тексты, написанные носителями других языков, значительно чаще ошибочно классифицируются как сгенерированные машиной, поскольку характеризуются более ограниченным лексическим разнообразием и более предсказуемыми синтаксическими конструкциями — то есть ровно теми признаками, на которые опирается детектор [3]. Для образовательной системы, в которой значительная часть академического письма выполняется на неродном для обучающегося языке, массовое внедрение таких инструментов означает не защиту академической честности, а воспроизводство дискриминации. В-третьих, надёжность детекции принципиально ограничена: незначительная редактура, перефразирование или последовательная обработка текста несколькими инструментами снижают её эффективность. Это следствие природы задачи, а не дефект конкретного продукта, и разрыв между генерацией и детекцией со временем скорее увеличивается [4]. Отсюда практический вывод: показатель детектора может рассматриваться исключительно как основание для дополнительного разбирательства — устного собеседования, защиты работы, проверки черновиков, — но не как самостоятельное доказательство нарушения. Институциональная политика, возлагающая на детектор доказательственную функцию, юридически уязвима и этически несостоятельна.

Три модели институционального реагирования. Анализ опубликованных университетских политик позволяет выделить три базовые модели, к которым сводится многообразие институциональных решений (таблица 1).

Таблица 1 — Модели институционального реагирования на генеративный ИИ

Модель	Содержание	Ограничения
Запретительная	Полный запрет использования генеративного ИИ; расчёт на средства детекции	Неисполнима на практике; переносит доказательственную нагрузку на ненадёжный инструмент; вытесняет использование ИИ в непрозрачную зону
Разрешительно-раскрывающая	Использование допускается при обязательном раскрытии характера и объёма применения	Требует внятных правил раскрытия и культуры доверия; сохраняет риск неполного декларирования
Интеграционная	Работа с ИИ рассматривается как формируемая компетенция и встраивается в результаты обучения и дизайн оценивания	Наиболее ресурсоёмка; требует переподготовки ППС и пересмотра оценочных средств

Запретительная модель при внешней строгости оказывается наименее эффективной: она не устраняет практику использования ИИ, а лишь делает её ненаблюдаемой, одновременно лишая преподавателя возможности обсуждать с обучающимся корректные способы работы с инструментом. Разрешительно-раскрывающая модель, отражённая, в частности, в принципах, согласованных ведущими исследовательскими университетами [5], исходит из противоположной посылки: обучающийся вправе использовать ИИ, но обязан заявить об этом, а образовательная организация обязана обучить его этому осознанно. Интеграционная модель развивает этот подход, предлагая дифференцировать допустимый уровень использования ИИ применительно к каждому оценочному заданию — от полного исключения до поощряемого применения; практическим воплощением такого подхода является шкала допустимого использования ИИ в оценивании [6]. Рекомендации ЮНЕСКО по использованию генеративного ИИ в образовании и научных исследованиях выстроены в той же логике, с акцентом на человекоцентричном подходе и институциональной политике, а не на запрете как таковом [7].

Дизайн оценивания как основной инструмент.

Если детекция не может служить опорой, единственным устойчивым механизмом остаётся пересмотр самих оценочных заданий. Уязвимым перед генеративным ИИ является задание, результат выполнения которого представляет собой текст, воспроизводящий известную информацию и не привязанный ни к личности обучающегося, ни к процессу его работы.

Устойчивость оценочного средства повышают следующие характеристики: опора на данные, собранные самим обучающимся; привязка к локальному, институциональному или профессиональному контексту, не представленному в обучающих корпусах моделей; включение устного компонента — защиты, собеседования, ответа на уточняющие вопросы; оценивание процесса, а не только итогового продукта (черновики, промежуточные версии, рефлексивный отчёт); требование критического анализа сгенерированного ИИ текста как части самого задания. Такой подход согласуется с концепцией безопасности оценивания, при которой защищённость достигается не наблюдением за нарушителем, а конструкцией задания, не дающей преимущества от обхода процедуры [8]. Существенно, что переработка оценочных средств лежит в зоне прямой ответственности вуза и не требует ни новых технологий, ни дополнительных лицензий.

Академическая честность как элемент внутренней системы обеспечения качества.

Стандарты и рекомендации для обеспечения качества в Европейском пространстве высшего образования устанавливают, что образовательные организации должны последовательно применять заранее опубликованные и единообразно реализуемые правила, охватывающие все этапы студенческого цикла, а оценивание должно осуществляться по объявленным критериям и процедурам [9]. Отсюда следует, что политика использования генеративного ИИ не является технической инструкцией: она входит в состав правил оценивания и, следовательно, подлежит опубликованию, последовательному применению и внутреннему мониторингу.

Практический смысл этого тезиса тройкий. Во-первых, наличие утверждённой, доведённой до сведения обучающихся и преподавателей политики использования ИИ становится проверяемым признаком зрелости внутренней системы обеспечения качества образования. Во-вторых, такая политика попадает в предмет внешней оценки: эксперты внешней экс-

пертной комиссии вправе рассматривать её наличие, непротиворечивость и фактическую реализуемость как индикатор соблюдения стандартов. В-третьих, это переводит проблему из индивидуальной плоскости («поймать обучающегося») в институциональную («обеспечить прозрачные и исполнимые правила»), что соответствует общей логике современного обеспечения качества.

В казахстанском контексте нормативная рамка академической честности задана законодательством об образовании и внутренними кодексами академической честности вузов, а инфраструктура проверки опирается преимущественно на системы выявления текстовых заимствований [указать источник — актуальная редакция закона и (или) документы конкретных вузов]. Существенным пробелом остаётся отсутствие единообразного подхода к раскрытию использования генеративного ИИ в письменных и квалификационных работах: при формальном запрете либо умолчании фактическое использование инструментов остаётся неурегулированным и ненаблюдаемым.

ЗАКЛЮЧЕНИЕ. Генеративный искусственный интеллект не создал проблему академической нечестности — он лишил работоспособности тот способ, которым система пыталась её решать. Инструментарий, ориентированный на поиск текстовых совпадений, утратил применимость там, где текст оригинален по форме и заимствован по существу. Попытка восстановить прежнюю логику средствами ИИ-детекции не выдерживает проверки: такие инструменты вероятностны, обходимы и систематически несправедливы по отношению к тем, кто пишет на неродном языке.

Устойчивый ответ предполагает, во-первых, переопределение академической честности через прозрачность процесса, а не через оригинальность продукта; во-вторых, пересмотр дизайна оценивания таким образом, чтобы обход процедуры не давал преимущества; в-третьих, закрепление политики использования ИИ в качестве полноценного элемента внутренней системы обеспечения качества образования, а значит — и предмета внешней оценки.

На основании изложенного организациям высшего и послевузовского образования Республики Казахстан могут быть предложены следующие рекомендации:

- принять на институциональном уровне политику использования генеративного ИИ, определяющую допустимые и недопустимые формы применения, порядок раскрытия и последствия несоблюдения, и включить её в состав правил оценивания;
- дифференцировать допустимый уровень использования ИИ по каждому оценочному заданию и указывать его в силлабусе;
- исключить использование показателей ИИ-детекторов в качестве самостоятельного доказательства нарушения, закрепив их статус как основания для дополнительного разбирательства;
- учесть риск систематического смещения детекторов против обучающихся, пишущих на неродном языке, что особенно значимо для многоязычной образовательной среды Казахстана;
- ввести обязательную декларацию об использовании ИИ в квалификационных и выпускных работах;
- пересмотреть оценочные средства с приоритетом устного компонента, локального контекста и оценивания процесса;
- включить наличие и реализуемость политики использования ИИ в состав показателей внутреннего мониторинга качества.

Практическая значимость результатов состоит в

возможности их использования при разработке институциональных политик и при проведении процедур внешней оценки качества. Направлением дальнейших исследований является эмпирический анализ факти-

ческих практик казахстанских вузов: наличия утверждённых политик использования ИИ, их содержания и степени реализуемости, а также восприятия этих правил обучающимися и преподавателями.

СПИСОК ЛИТЕРАТУРЫ:

1. The Fundamental Values of Academic Integrity. — 3rd ed. — [S. l.]: International Center for Academic Integrity, 2021. — 12 p. [свернуть выходные данные и URL]
2. Authorship and AI tools: COPE position statement [Электронный ресурс] / Committee on Publication Ethics. — 2023.
3. GPT detectors are biased against non-native English writers / W. Liang, M. Yuksekgonul, Y. Mao [et al.] // Patterns. — 2023. — Vol. 4, No. 7.
4. Can AI-Generated Text be Reliably Detected? / V. S. Sadasivan, A. Kumar, S. Balasubramanian [et al.] // arXiv preprint. — 2023. — arXiv:2303.11156.
5. Russell Group principles on the use of generative AI tools in education [Электронный ресурс] / Russell Group. — 2023.
6. The AI Assessment Scale (AIAS): A framework for ethical integration of generative AI in educational assessment / M. Perkins, L. Furze, J. Roe, J. MacVaugh // Journal of University Teaching and Learning Practice. — 2024.
7. Guidance for generative AI in education and research. — Paris: UNESCO, 2023.
8. Dawson P. Defending Assessment Security in a Digital World: Preventing E-Cheating and Supporting Academic Integrity in Higher Education / P. Dawson. — London: Routledge, 2021.
9. Standards and Guidelines for Quality Assurance in the European Higher Education Area (ESG). — Brussels: EURASHE, 2015. — 32 p.
10. Закон Республики Казахстан «Об образовании»

REFERENCES:

1. The Fundamental Values of Academic Integrity (2021). 3rd ed. International Center for Academic Integrity.
2. Authorship and AI tools: COPE position statement (2023). Committee on Publication Ethics.
3. Liang, W., Yuksekgonul, M., Mao, Y. et al. (2023). GPT detectors are biased against non-native English writers. Patterns, 4(7)
4. Sadasivan, V. S., Kumar, A., Balasubramanian, S. et al. (2023). Can AI-Generated Text be Reliably Detected? arXiv preprint arXiv:2303.11156.
5. Russell Group principles on the use of generative AI tools in education (2023). Russell Group.
6. Perkins, M., Furze, L., Roe, J., MacVaugh, J. (2024). The AI Assessment Scale (AIAS): A framework for ethical integration of generative AI in educational assessment. Journal of University Teaching and Learning Practice
7. Guidance for generative AI in education and research (2023). Paris: UNESCO.
8. Dawson, P. (2021). Defending Assessment Security in a Digital World: Preventing E-Cheating and Supporting Academic Integrity in Higher Education. London: Routledge.
9. Standards and Guidelines for Quality Assurance in the European Higher Education Area (ESG) (2015). Brussels: EURASHE
10. Zakon Respubliki Kazakhstan "Ob obrazovanii" [Law of the Republic of Kazakhstan "On Education"]

AUTHOR INFORMATION:

Конусова Перизат — магистр технических наук, руководитель информационно-аналитических проектов, Независимое агентство аккредитации и рейтинга (г. Астана, Республика Казахстан), e-mail: perizat.konusova@mail.ru

Фартунова Мария Атанасова — доктор PhD, доцент, Софийский университет им. Св. Климента Охридского (г. София, Республика Болгария), e-mail: m.fartunova@abv.bg

Конусова Перизат — техника ғылымдарының магистрі, ақпараттық-аналитикалық жобалардың жетекшісі, Аккредиттеу және рейтинг тәуелсіз агенттігі (Астана қ., Қазақстан Республикасы), e-mail: perizat.konusova@mail.ru

Фартунова Мария Атанасова — PhD докторы, доцент, Әулие Климент Охридский атындағы София университеті (София қ., Болгария), e-mail: m.fartunova@abv.bg

Konussova Perizat — Master of Technical Sciences, Head of Information and Analytical Projects, Independent Agency for Accreditation and Rating (Astana, Republic of Kazakhstan), e-mail: perizat.konusova@mail.ru

Fartunova Maria Atanasova — PhD, Associate Professor, Sofia University "St. Kliment Ohridski" (Sofia, Republic of Bulgaria), e-mail: m.fartunova@abv.bg

МРНТИ 06.00.32
УДК 338.2:658.155(574)

DOI 10.58319/26170493_2026_2_45

ИНТЕГРАЦИЯ ESG-ФАКТОРОВ В СТРАТЕГИЧЕСКОЕ УПРАВЛЕНИЕ НЕФТЕГАЗОВЫМИ КОМПАНИЯМИ В УСЛОВИЯХ ГЛОБАЛЬНЫХ ЭНЕРГЕТИЧЕСКИХ ИЗМЕНЕНИЙ

Урлова А.Ш.¹, Саймагамбетова Г.А.¹, Увакбаева Г.Б.², Есбергенова Б.М.¹

¹Каспийский государственный университет технологий и инжиниринга имени Ш. Есенова, Актау, Казахстан;

²Казахский национальный университет имени аль-Фараби, Алматы, Казахстан

АННОТАЦИЯ. В статье рассматриваются актуальные вопросы влияния ESG-факторов (экологических, социальных и управленческих) на стратегии развития нефтегазовых компаний Казахстана в условиях глобального энергоперехода и усиления требований устойчивого развития. Особое внимание уделяется тому, как международные стандарты ESG, изменение критериев инвестирования и ужесточение экологических норм трансформируют корпоративные приоритеты, операционные процессы и долгосрочные планы предприятий отрасли. Анализируется влияние таких факторов, как декарбонизация, энергоэффективность, социальная ответственность и прозрачность корпоративного управления, на конкурентоспособность компаний на международных рынках.

Цель исследования – определить степень и механизмы влияния ESG-подходов на стратегическое развитие нефтегазовых компаний Казахстана, а также выявить ключевые направления адаптации отрасли к новым требованиям устойчивого развития.

Методологический подход включает исследование внутренних и внешних особенностей, а также комплексный анализ стандартов ESG, оценку корпоративной отчётности нефтегазовых предприятий, стратегический и SWOT-анализ. Это позволило определить векторы ESG-изменений. Задачами статьи является обоснование влияния ESG на стратегию развития нефтегазового кластера с учетом национальных особенностей и определение практических подходов, направленных