

УДК 37:004.8

Для цитирования: Бжекшиев А.Р. Проблема оценивания в условиях распространения генеративного искусственного интеллекта: авторство работы как условие интерпретации результата // Научно-методический журнал «Поиск научных решений». 2026. № 1. С. 81–90.

DOI: 10.61077/2949-4818-2026-1-81-90

**Проблема оценивания в условиях распространения
генеративного искусственного интеллекта: авторство работы
как условие интерпретации результата**

Бжекшиев Анзор Русланович

Государственное бюджетное учреждение
дополнительного профессионального образования «Центр непрерывного
повышения профессионального мастерства педагогических работников»
Министерства просвещения и науки Кабардино-Балкарской Республики

Аннотация. Статья представляет концептуальный анализ проблемы оценивания в условиях массовой доступности генеративных языковых моделей. Автор обосновывает тезис о том, что данная проблема является артефактом специфической архитектуры оценивания, в которой отчуждаемый результат учащегося (текст, решение задачи, проект) наделяется статусом основного эмпирического основания педагогической интерпретации. В исследовании характеризуется спектр адаптивных стратегий образовательных организаций и выявляется их общая направленность на сохранение существующей архитектуры и восстановление достоверности предположения об авторстве работы. Анализ показывает, что при организации оценивания через непосредственно наблюдаемую деятельность ученика зависимость интерпретации от предположения об авторстве существенно снижается, однако масштабирование таких форм ограничено организационными условиями массового образования.

Ключевые слова: оценка, генеративный искусственный интеллект, автор, архитектура, оценка, интерпретация, результат, валидность, образование.

**The problem of assessment in the context of the spread
of generative artificial intelligence: authorship of a work as a condition
for interpreting the result**

Bzhekshiev Anzor Ruslanovich

State Budgetary Institution of Additional Professional Education
"Center for Continuous Development of Professional Mastery of Teaching Staff"
of the Ministry of Enlightenment and Science of the Kabardino-Balkarian Republic

Abstract. The article presents a conceptual analysis of the problem of assessment in the context of mass availability of generative language models. The author substantiates the thesis that this problem is an artifact of a specific assessment architecture in which the alienated result of the student (text, problem solution, project) is given the status of the main empirical basis of pedagogical interpretation. The study characterizes the range of adaptive strategies of educational organizations and identifies their general focus on preserving the existing architecture and restoring the credibility of the assumption of authorship of the work. The analysis shows that when organizing assessment through the student's directly observed activity, the dependence of interpretation on the assumption of authorship is significantly reduced, but the scale of such forms is limited by the organizational conditions of mass education.

Keywords: assessment, generative artificial intelligence, author, architecture, evaluation, interpretation, result, validity, education.

Генеративные системы искусственного интеллекта, способные порождать связные тексты, решать задачи и выстраивать аргументацию, стали массово доступны для участников образовательного процесса. Профессиональное сообщество реагирует на эту ситуацию разнообразными мерами [Ивахненко, Никольский 2023]: от запретов на использование и внедрения детекторов ИИ-генерированного текста до изменения форматов заданий и требований фиксировать процесс работы. Предлагаемые меры, однако, не привели к формированию устойчивого подхода, применимого в условиях массового образования, а по мере развития языковых моделей перспективы их эффективности остаются неопределёнными. В этих условиях предметом анализа становится сам способ организации оценивания, в пределах которого затруднения, связанные с использованием ИИ, воспроизводятся.

Основной корпус современно исследовательской литературы, посвященной проблемам оценивания в условиях генеративного ИИ, сложился преимущественно в поле высшего образования. Именно университетская среда впервые столкнулась с массовым использованием языковых моделей студентами и начала документировать его последствия. Вместе с тем описываемая в этих исследованиях зависимость оценочного вывода от предъявленного продукта носит более общий характер и в полной мере

распространяется на условия массовой школы, где ограниченность ресурсов и масштаб процесса делают эту зависимость более выраженной.

Предпринимаемые меры при всем их разнообразии различаются прежде всего тем, на каком этапе оценочной процедуры действуют. Первая группа мер связана с попыткой установить, является ли предъявленная работа результатом самостоятельной деятельности ученика, то есть представляет собой попытку распознать факт использования ИИ уже после того, как работа сдана. Сюда относятся как специализированные программные инструменты (Turnitin AI Writing Detection, GPTZero, Copyleaks), так и попытки педагогов самостоятельно опознать ИИ-генерированный текст по стилистическим признакам. Масштабное тестирование детекторов, проведенное Weber-Wulff и коллегами [Weber-Wulff et al. 2023], показало, что ни один из доступных инструментов не обеспечивал надежного различения человеческого и ИИ-генерированного текста: инструменты демонстрировали склонность классифицировать проверяемый материал как написанный человеком и существенно теряли в эффективности при малейшем редактировании или парафразе. Liang и соавторы [Liang et al. 2023] установили, что детекторы систематически ошибаются в отношении текстов, написанных носителями языка, квалифицируя подлинные студенческие работы как сгенерированные искусственным интеллектом. Это ограничение не является исключительно техническим: тексты носителей обладают меньшей лексической вариативностью, которую алгоритмы интерпретируют как признак машинной генерации, и устранение этого эффекта требует иного принципа различения. Показательно, OpenAI, разработчик ChatGPT, прекратил поддержку собственного классификатора по причине недостаточной точности.

Что касается распознавания со стороны педагогов, Hardie и коллеги [Hardie et al. 2024] установили, что специальная подготовка преподавателей по распознаванию ИИ-текстов повышала их способность обнаруживать генерированные работы, но одновременно приводила к значительному росту ложноположительных результатов – преподаватели начинали видеть признаки ИИ в подлинных студенческих работах.

Следующая группа мер направлена на изменение самих условий выполнения задания таким образом, чтобы использование ИИ стало невозможным или бесполезным. К этой группе относятся перенос оценочных процедур в контролируемую среду (аудиторные работы без доступа к устройствам, прокторинг), использование устных форматов в функции верификации самостоятельности выполнения (защита проекта, устный экзамен, собеседование), а также проектирование заданий, предполагающих обращение

к уникальному личному опыту ученика. Однако аудиторный контроль при регулярном оценивании по множеству предметов создает организационную нагрузку, несовместимую с реальными условиями работы массовой школы, а устные форматы требуют индивидуальной работы с каждым учащимся, возможности для которой у педагога, работающего с несколькими классами, как правило, нет. В то же время, проектирование заданий, устойчивых к ИИ, наталкивается на эмпирическое ограничение: Borges и соавторы [Borges et al. 2024] продемонстрировали, что ChatGPT успешно справляется с экзаменационными вопросами по 50 университетским курсам, а Hardie и коллеги [Hardie et al. 2024] установили, что из 17 протестированных типов заданий большинство оказалось уязвимым для генеративного ИИ. Kofinas, Tsay и Pike [Kofinas, Tsay, Pike 2025] убедительно показали, что даже так называемые аутентичные задания не обеспечивают гарантированного различия: эксперты в их исследовании одновременно пропускали ИИ-генерированные ответы и ошибочно квалифицировали подлинные студенческие работы.

Последняя группа мер связана с требованием зафиксировать процесс работы через промежуточные черновики, журналы работы, логирование обращений к ИИ или поэтапную сдачу. Логика состоит в том, что видимый процесс позволяет судить о самостоятельности результата. Помимо значительной дополнительной нагрузки на педагога, который должен в этом случае анализировать серию промежуточных версий, фиксация процесса не устраняет зависимости оценочного вывода от интерпретации предъявленного материала: промежуточные черновики так же являются продуктами, которые могут быть сгенерированы в несколько итераций. Журнал работы может быть составлен ретроспективно, а педагог, анализируя цепочку черновиков, по-прежнему решает тот же вопрос о происхождении предъявленного, но теперь в отношении каждого отдельного элемента цепочки.

Таблица 1.

Меры противодействия
использованию генеративного ИИ в оценивании

Группа мер	Объект воздействия	Логика действия	Ограничение
Детекция ИИ-генерированного текста	Итоговый продукт (после сдачи)	Реконструкция происхождения текста по стилистическим и статистическим признакам	Недостижимая надежность различения; систематические ошибки в отношении текстов не носителей языка

Изменение условий выполнения задания	Среда выполнения	Исключение доступа к ИИ или проектирование заданий, устойчивых к генерации	Организационная нагрузка, несовместимая с массовой школой; эмпирическая уязвимость заданий
Фиксация процесса работы	Промежуточные версии продукта	Восстановление видимости процесса через серию объектов	Мультипликация того же вопроса о происхождении в отношении каждого элемента цепочки

Все группы мер обнаруживают общее ограничение. Детекция работает с итоговым текстом и пытается реконструировать его происхождение. Изменение условий стремится исключить определенный способ создания текста. Фиксация процесса множит количество объектов, в отношении которых ставится тот же вопрос о происхождении. Во всех случаях оценочный вывод строится на анализе предъявленного учеником материала, и его устойчивость определяется тем, насколько достоверно можно установить, как именно этот материал был произведен.

То, что объединяет все рассмотренные меры и задает границы их эффективности, может быть описано через устройство самой процедуры оценочного вывода. Оценивание, каким оно сложилось в массовом образовании в качестве доминирующей практики, представляет собой цепочку, в которой ученик предъявляет результат своей работы – текст, решение задачи, проект, ответ на вопрос, – педагог интерпретирует этот результат как свидетельство определенного уровня освоения содержания, и на основании этой интерпретации принимает решение – выставляет отметку, формулирует обратную связь, допускает к следующему этапу обучения. Эта цепочка – продукт, интерпретация, решение – образует продуктоцентричную архитектуру оценивания, в рамках которой предъявленный учеником продукт служит основным эмпирическим основанием интерпретации.

Устойчивость интерпретации в рамках этой архитектуры зависит от предположения о происхождении продукта. Когда педагог ставит отметку за сочинение, он исходит из того, что текст порожден деятельностью ученика – его пониманием темы, умением выстраивать аргументацию и работать с источниками. Если это предположение не выполняется, интерпретация утрачивает основание: тот же текст, произведенный другим субъектом, уже не может служить свидетельством освоения содержания.

Такого рода зависимость не нова: готовые домашние задания, списывание, помощь со стороны других лиц и шпаргалки нарушали предположение о происхождении продукта в пределах той же оценочной архитектуры [Шмелева 2016]. Dawson, Bearman, Dollinger и Boud [Dawson et al.

2024], анализируя эту проблему, показали, что она, традиционно обсуждаемая в категориях академической честности, по существу является проблемой валидности оценивания в том смысле, который придает этому понятию [Messick 1989].

Генеративный ИИ в рамках этой архитектуры выступает как инструмент, способный порождать уникальный продукт, неотличимый от продукта деятельности учащегося и не оставляющий следов источника. В отличие от прежних форм подмены авторства результат генерации не поддается верификации средствами анализа самого продукта, и предположение о происхождении, на котором держится цепочка интерпретации, утрачивает опору.



Рис. 1. Архитектура оценочного вывода и точка уязвимости

Между тем процедура оценивания допускает и иную организацию, при которой интерпретация опирается на непосредственно наблюдаемую деятельность ученика [Black, Wiliam 1998; Паскова 2024]. Те же форматы, которые относятся к мерам контроля – устный опрос, защита проекта, практическая работа, – приобретают иную функцию, если используются не для верификации продукта, а как самостоятельное основание интерпретации. Устный опрос, в ходе которого ученик ориентируется в содержании в реальном времени, практическая работа, выполняемая совместно с педагогом, защита проекта, предполагающая не воспроизведение заранее подготовленного текста, а свободное владение материалом – во всех этих случаях основанием

интерпретации служит наблюдаемая способность ученика действовать с содержанием, а не зафиксированный результат этой деятельности.

В таких конфигурациях искусственный интеллект не исчезает из учебной среды. Он остается доступным так же, как учебник, справочник, совет одноклассника или любой другой источник. Однако его использование перестает быть решающим для оценочного вывода, поскольку интерпретация опирается на то, как ученик ориентируется в содержании здесь и сейчас, в ситуации, которую педагог наблюдает непосредственно.

Несмотря на то что устный ответ и защита проекта также могут быть подготовлены с использованием ИИ, данное обстоятельство не устраняет структурного различия между оцениванием по продукту и оцениванием по наблюдаемой деятельности. В продуктоцентричной модели педагог работает с материалом, понимание происхождения которого ему недоступно. В модели, основанной на наблюдении, педагог видит, как ученик работает с содержанием, и интерпретация строится на этой наблюдаемой способности. Ученик, который подготовился с помощью ИИ и в результате способен свободно ориентироваться в материале, отвечать на вопросы и выстраивать рассуждение в реальном времени, демонстрирует ту компетенцию, которую оценивание призвано зафиксировать. Ученик, который не способен этого сделать, обнаруживает пределы своего освоения содержания вне зависимости от того, какие инструменты он использовал при подготовке. Разумеется, если предметом оценивания является именно умение создавать текст определённого типа, наблюдаемая деятельность не замещает эту компетенцию — она устойчива не ко всем задачам оценивания, а к проблеме, порождаемой зависимостью интерпретации от предположения об авторстве. Регулярное применение подобных форм оценивания, однако, требует организационных условий, которыми массовая школа, как правило, не располагает.

В этом контексте продуктоцентричная архитектура оценивания, существующая как доминирующая конфигурация массового образования, не является ни ошибкой, ни результатом педагогической или управленческой некомпетентности. Она представляет собой рациональный ответ системы на условия массового обучения: ограниченное время урока, необходимость документируемого и сопоставимого результата, а также требования формализованных оценочных процедур. В то время как организация оценивания через наблюдаемую деятельность предполагает более длительное и индивидуализированное взаимодействие с учеником, что делает ее регулярное применение в массовом режиме затруднительным. Bearman, Tai, Dawson, Boud и Ajjawī [Bearman et al. 2024] указывают в этой связи на то, что изолированные

изменения отдельных заданий не решают проблемы без пересмотра организационных условий в целом.

С учетом этих ограничений аргумент настоящей статьи является сознательно узким. Статья не предлагает реформы системы оценивания и не утверждает, что продуктоцентричная модель должна быть заменена какой-либо альтернативой. Она фиксирует, что проблема оценивания в условиях генеративного ИИ является следствием конкретной архитектуры оценочного вывода и что при иной организации процесса значимость предположения о происхождении продукта для интерпретации существенно снижается или исчезает. Вопрос о том, при каких организационных, нормативных и ресурсных условиях такие конфигурации могут быть масштабированы, представляет собой самостоятельный предмет исследования.

Проведенный анализ позволяет уточнить направление дальнейшего обсуждения проблемы оценивания. Совершенствование мер, ориентированных на восстановление достоверности предположения о происхождении предъявленного продукта, не приводит к устойчивому решению, поскольку эти меры сохраняют ту архитектуру оценивания, в пределах которой проблема воспроизводится. При сохранении прежней оптики усилия профессионального сообщества сосредотачиваются на мерах контроля, не способных принести устойчивый результат, тогда как достоверность оценочного вывода – его способность служить свидетельством освоения содержания – остается под вопросом.

В этой связи аналитически и практически значимым становится поиск таких форм организации оценивания, при которых предположение о происхождении продукта утрачивает определяющую роль в интерпретации результата. Разработка подобных форм и исследование условий их масштабирования составляют наиболее значимую задачу для продолжения этой дискуссии.

Список источников и литературы

1. 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.
2. Паскова А.А. Возможности интеграции технологий генеративного искусственного интеллекта в процессы формирующего оценивания в высшем образовании // Вестник Майкопского государственного технологического университета. 2024. Т. 16. № 2. С. 98–109.

3. Шмелева Е.Д. Плагиат и списывание в российских вузах: роль образовательной среды и индивидуальных характеристик студента // Вопросы образования. 2016. № 1. С. 84–109.
4. Bearman M., Ajjawi R. Learning to Work with the Black Box: Pedagogy for a World with Artificial Intelligence // British Journal of Educational Technology. 2023. Vol. 54. No. 5. P. 1160–1173.
5. Bearman M., Tai J., Dawson P., Boud D., Ajjawi R. Developing Evaluative Judgement for a Time of Generative Artificial Intelligence // Assessment & Evaluation in Higher Education. 2024. Vol. 49. No. 6. P. 893–905.
6. Black P., Wiliam D. Assessment and Classroom Learning // Assessment in Education: Principles, Policy & Practice. 1998. Vol. 5. No. 1. P. 7–74.
7. Borges B., Foroutan N., Bayazit D., Sotnikova A., Bosselut A., EPFL Grader Consortium, EPFL Data Consortium. Could ChatGPT Get an Engineering Degree? Evaluating Higher Education Vulnerability to AI Assistants // Proceedings of the National Academy of Sciences. 2024. Vol. 121. No. 49. e2414955121.
8. Dawson P., Bearman M., Dollinger M., Boud D. Validity Matters More than Cheating // Assessment & Evaluation in Higher Education. 2024. Vol. 49. No. 7. P. 1005–1016.
9. Hardie L., Lowe J., Pride M., Waugh K., Hauck M., Ryan F., Gooch D., Patent V., McCartney K., Maguire C., Richards M., Richardson H. Developing Robust Assessment in the Light of Generative AI Developments. NCFE; The Open University, 2024.
10. Kofinas A.K., Tsay C.H.-H., Pike D. The Impact of Generative AI on Academic Integrity of Authentic Assessments within a Higher Education Context // British Journal of Educational Technology. 2025. Vol. 56. No. 6. P. 2522–2549.
11. Liang W., Yuksekgonul M., Mao Y., Wu E., Zou J. GPT Detectors Are Biased against Non-Native English Writers // Patterns. 2023. Vol. 4. No. 7. 100779.
12. Lodge J., Howard S., Bearman M., Dawson P. Assessment Reform for the Age of Artificial Intelligence. Tertiary Education Quality and Standards Agency (TEQSA), 2023.
13. Messick S. Validity // Educational Measurement / ed. R.L. Linn. 3rd ed. American Council on Education / Macmillan, 1989. P. 13–103.
14. Weber-Wulff D., Anohina-Naumeca A., Bjelobaba S., Foltýnek T., Guerrero-Dib J., Popoola O., Šigut P., Waddington L. Testing of Detection Tools for AI-Generated Text // International Journal for Educational Integrity. 2023. Vol. 19. No. 1. 26.

References

1. Ivakhnenko, E.N., Nikolsky, V.S. ChatGPT in Higher Education and Science: Threat or Valuable Resource? *Higher Education in Russia*, 2023, Vol. 32, no. 4, pp. 9–22, in Russian.

2. Paskova, A.A. Possibilities of Integrating Generative Artificial Intelligence Technologies into Formative Assessment Processes in Higher Education, *Bulletin of Maikop State Technological University*, 2024, vol. 16, no. 2, pp. 98–109, *in Russian*.
3. Shmeleva, E.D. Plagiarism and Cheating in Russian Universities: The Role of Educational Environment and Student Individual Characteristics, *Educational Studies*, 2016, no. 1, pp. 84–109, *in Russian*.
4. Bearman, M., Ajjawi, R. *Learning to Work with the Black Box: Pedagogy for a World with Artificial Intelligence* // *British Journal of Educational Technology*. 2023. Vol. 54. No. 5. Pp. 1160–1173.
5. Bearman, M., Tai, J., Dawson, P., Boud, D., Ajjawi, R. *Developing Evaluative Judgement for a Time of Generative Artificial Intelligence*, *Assessment & Evaluation in Higher Education*, 2024, vol. 49, no. 6, pp. 893–905.
6. Black, P., Wiliam, D. *Assessment and Classroom Learnin*, *Assessment in Education: Principles, Policy & Practice*, 1998 , vol. 5, no. 1, pp. 7–74.
7. Borges, B., Foroutan, N., Bayazit, D., Sotnikova, A., Bosselut, A., EPFL Grader Consortium, EPFL Data Consortium. *Could ChatGPT Get an Engineering Degree? Evaluating Higher Education Vulnerability to AI Assistants*, *Proceedings of the National Academy of Sciences*, 2024, vol. 121, no. 49. e2414955121.
8. Dawson, P., Bearman, M., Dollinger, M., Boud, D. *Validity Matters More than Cheating* // *Assessment & Evaluation in Higher Education*, 2024, vol. 49, no. 7, pp. 1005–1016.
9. Hardie, L., Lowe, J., Pride, M., Waugh, K., Hauck, M., Ryan, F., Gooch, D., Patent, V., McCartney, K., Maguire, C., Richards, M., Richardson, H. *Developing Robust Assessment in the Light of Generative AI Developments*. NCFE; The Open University, 2024.
10. Kofinas, A.K., Tsay, C.H.-H., Pike, D. *The Impact of Generative AI on Academic Integrity of Authentic Assessments within a Higher Education Context*, *British Journal of Educational Technology*, 2025, vol. 56, no. 6., pp. 2522–2549.
11. Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., Zou, J. *GPT Detectors Are Biased against Non-Native English Writers*, *Patterns*. 2023, vol. 4, no. 7. 100779.
12. Lodge, J., Howard, S., Bearman, M., Dawson, P. *Assessment Reform for the Age of Artificial Intelligence*. Tertiary Education Quality and Standards Agency (TEQSA), 2023.
13. Messick, S. *Validity*, *Educational Measurement*. Ed. R.L. Linn. 3rd ed. American Council on Education / Macmillan, 1989. Pp. 13–103.
14. Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P., Waddington, L. *Testing of Detection Tools for AI-Generated Text*, *International Journal*