

SUN'IY INTELEKTNING TA'LIM SIFATINI BAHOLASHDAGI ISHONCHLILIGI VA XOLISLIGI MUAMMOLARI

Dilmurodov Shaxboz Sirojiddin o'g'li

Muhammad al-Xorazmiy nomidagi Toshkent

axborot texnologiyalari universiteti talabasi

E-mail: dilmurodovshaxboz94@gmail.com

Annotatsiya Ushbu maqolada sun'iy intellektning ta'lim sifatini baholashdagi ishonchliligi va xolisligi muammolari tahlil qilinadi. Avtomatlashtirilgan baholash tizimlari va generativ SI vositalarining samaradorligi ko'rib chiqilib, algoritmik biasning asosiy sabablari (trening ma'lumotlaridagi nomutanosiblik, madaniy-lingvistik farqlar, "qora quti" effekti) aniqlangan. Tadqiqot natijalari shuni ko'rsatadiki, SI baholashda xatolik va noxolislik ehtimoli yuqori bo'lib, ayniqsa ozchilik guruhlar uchun adolatsizlik kuzatiladi. Muammolarni yumshatish uchun ma'lumotlar bazasini diversifikatsiya qilish, algoritmlarni audit qilish, fairness metrikalarini qo'llash va inson nazoratini joriy etish taklif etiladi. Maqola ta'limda SI ni adolatli qo'llash bo'yicha amaliy tavsiyalar beradi.

Kalit so'zlar: sun'iy intellekt, ta'lim baholash, ishonchlilik, xolislik, algoritmik bias, fairness, shaffoflik.

Аннотация: В статье анализируются проблемы надёжности и объективности искусственного интеллекта (ИИ) при оценке качества образования. Рассматривается эффективность систем автоматизированного оценивания и инструментов генеративного ИИ, а также основные причины алгоритмического bias: несбалансированность обучающих данных, культурно-языковые различия, эффект «чёрного ящика» и недостатки дизайна алгоритмов. Результаты показывают, что использование ИИ в оценивании сопряжено с высоким риском ошибок и несправедливости, особенно в отношении групп. Для снижения этих рисков предлагается диверсификация баз данных, регулярный аудит алгоритмов, применение метрик справедливости, повышение прозрачности (explainable AI) и обязательный человеческий контроль. Статья даёт практические рекомендации по справедливому и эффективному применению ИИ в образовании.

Ключевые слова: искусственный интеллект, оценка в образовании, надёжность, объективность, алгоритмический bias, прозрачность.

Abstract: This article analyzes the issues of reliability and fairness in the use of artificial intelligence (AI) for assessing the quality of education. The effectiveness of automated assessment systems and generative AI tools is examined, and the main causes of algorithmic bias are identified, including imbalances in training data, cultural-linguistic differences, and the "black box" effect. The study findings indicate

that AI-based assessment carries a high probability of errors and unfairness, particularly leading to disadvantages for minority and underrepresented groups. To mitigate these issues, the following measures are proposed: diversifying training datasets, conducting regular algorithm audits, applying fairness metrics, and introducing mandatory human oversight. The article offers practical recommendations for the fair and effective application of AI in education.

Keywords: artificial intelligence, educational assessment, reliability, fairness, algorithmic bias, fairness metrics, transparency.

Kirish: Zamonaviy ta'limda sun'iy intellekt (SI) texnologiyalari, xususan avtomatlashtirilgan baholash tizimlari va generativ AI vositalari tez sur'atlar bilan qo'llanilmoqda. Ular baholash jarayonini tezlashtirish, izchillikni oshirish va shaxsiylashtirilgan yondashuvni ta'minlash orqali katta imkoniyatlar beradi. Biroq, SI ning ta'lim sifatini baholashdagi qo'llanilishi jiddiy muammolarni ham keltirib chiqarmoqda. Eng dolzarblari — ishonchlilik (reliability) va xolislik (fairness) masalalari bo'lib, trening ma'lumotlaridagi nomutanosiblik, madaniy-lingvistik farqlar va "qora quti" effekti tufayli algoritmik bias yuzaga kelib, ozchilik guruhlarini (boshqa tilli talabalar, madaniy farqli yoki kam ta'minlangan talabalar) uchun sistematik adolatsizliklar paydo bo'lmoqda. So'nggi tadqiqotlar (2023–2025) shuni ko'rsatadiki, AI baholash tizimlari inson bahosiga nisbatan yuqori konsistensiyaga ega bo'lsa-da, fairness (demographic parity, equal opportunity) buzilishi va madaniy biaslar hali ham dolzarb muammo sifatida qolmoqda. Bunday holatlar ta'limdagi mavjud tengsizliklarni kuchaytirishi mumkin. Ushbu maqola sun'iy intellektning ta'lim baholashidagi ishonchliligi va xolisligi muammolarini tahlil qilishga bag'ishlangan. Algoritmik biasning asosiy sabablari aniqlanib, ularni yumshatish yo'llari (ma'lumotlar diversifikatsiyasi, audit, fairness metrikalari, inson nazorati) ko'rib chiqiladi. Maqsad — ta'limda SI ni adolatli qo'llash bo'yicha amaliy tavsiyalar berish, shu jumladan O'zbekiston ta'lim tizimining xususiyatlarini hisobga olgan holda.

Methodology (Adabiyotlar tahlili va metodlar): Ushbu tadqiqot nazariy-tahliliy xarakterga ega bo'lib, sun'iy intellektning ta'lim baholashidagi ishonchliligi va xolisligi muammolarini mavjud ilmiy adabiyotlar asosida chuqur o'rganishga qaratilgan. Adabiyotlar tahlili 2023–2026 yillarda chop etilgan ilmiy nashrlarga asoslanadi. Scopus, Google Scholar, arXiv, ResearchGate va boshqa bazalardan "AI educational assessment fairness", "algorithmic bias in automated grading", "bias mitigation in educational AI", "fairness in AI grading" kalit so'zlari bo'yicha 40 dan ortiq peer-reviewed maqola, konferensiya materiallari va rasmiy hisobotlar (UNESCO, U.S. Department of Education, European Commission) tanlab olindi va tahlil qilindi. Asosiy e'tibor algoritmik biasning sabablari (trening ma'lumotlaridagi nomutanosiblik, madaniy-lingvistik farqlar, "qora quti" effekti), fairness metrikalari

(demographic parity, equalized odds, equal opportunity) va yumshatish strategiyalariga (pre-processing, in-processing, post-processing) qaratildi. O‘zbekiston va rivojlanayotgan mamlakatlar ta’lim tizimlarida til ko‘pligi va madaniy xilma-xillik tufayli biasning kuchayishi masalasi ham ko‘rib chiqildi.

Tadqiqot metodlari quyidagilardan iborat: sistematik adabiyot sharhi (systematic literature review) – tanlangan manbalarni dolzarblik, ilmiy sifat va ta’lim kontekstiga oidlik mezonlari asosida saralash; kontent tahlili – bias sabablari, fairness metrikalari va mitigation usullarini tematik guruhlash va tasniflash; taqqoslama tahlili – turli tadqiqotlarda qo‘llanilgan yondashuvlarning samaradorligi va cheklovlarini solishtirish; kontekstual moslashtirish – olingan natijalarni O‘zbekiston ta’lim tizimining xususiyatlariga (ko‘p tillilik, resurs cheklovlari, madaniy farqlar) moslashtirish.

Natijalar: Tadqiqot natijalari shuni ko‘rsatadiki, sun’iy intellekt (SI) asosidagi ta’lim baholash tizimlari ishonchlilik (reliability) jihatidan inson bahosiga nisbatan yuqori konsistensiyaga ega: baholar barqarorroq va insoniy sub’ektiv omillardan (charchoq, kayfiyat) kamroq ta’sirlanadi. Biroq, xolislik (fairness) masalasida jiddiy kamchiliklar mavjud. Algoritmik biasning asosiy sabablari quyidagilar:

- Trening ma’lumotlaridagi nomutanosiblik va madaniy-lingvistik cheklovlar (non-native speakerlar, ozchilik guruhlari uchun sistematik past baholar).
- “Qora quti” effekti (shaffoflik yetishmasligi).
- Insoniy biasning AI ga o‘tkazilishi.

So‘nggi tadqiqotlar (2023–2026) shuni tasdiqlaydiki, AI baholashda demographic parity va equal opportunity metrikalari buzilishi kuzatiladi, ayniqsa til farqli, madaniy ozchilik yoki kam ta’minlangan talabalar uchun adolatsizlik yuqori (masalan, non-native English speakerlar esselari past baholanadi).

Quyidagi jadvalda fairness metrikalari va ularning ta’lim baholashidagi qo‘llanilishi keltirilgan:

Fairness metriki	Ta’rifi	Ta’lim kontekstida misol	Natija (umumiy)
Demographic Parity	Pozitiv natijalar (masalan, yuqori baho) guruhlar bo’yicha teng bo’lishi kerak	AI gradingda barcha guruhlar uchun o’xshash baho foizi	Ko‘pincha buziladi (ozchilik guruhlari pastroq)
Equal Opportunity	True Positive Rate (to‘g‘ri yuqori baholanganlar) guruhlar bo’yicha teng	Haqiqiy yaxshi talabalar teng baholanadi	Buzilishi kuzatiladi
Equalized Odds	True Positive va False Positive Rate teng bo’lishi	Xato baholar (past baholash) guruhlar bo’yicha teng	Eng ko‘p muammo shu metrikda

Biasni kamaytirish choralari samarali bo'lishi uchun ma'lumotlar bazasini diversifikatsiya qilish, muntazam auditlar o'tkazish, fairness metrikalarini qo'llash (masalan, threshold adjustment yoki adversarial debiasing), shuningdek, explainable AI va human-in-the-loop yondashuvlari muhimdir. O'zbekiston kontekstida esa til ko'pligi va madaniy xilma-xillik tufayli bias kuchayishi mumkin, shuning uchun mahalliy representativ ma'lumotlar bilan trening va kontekstual moslashtirish alohida ahamiyatga ega.

Muhokama: Natijalar sun'iy intellektning ta'lim baholashidagi ikki tomonlama xususiyatini ko'rsatadi: SI tizimlari baholash jarayonini tez va izchil qiladi, insoniy sub'ektivlikni kamaytiradi, lekin algoritmik bias tufayli xolislik muammosi saqlanib qolmoqda. AI yuqori konsistensiyaga ega bo'lsa-da, trening ma'lumotlaridagi nomutanosiblik, madaniy-lingvistik farqlar va "qora quti" effekti ozchilik guruhlari (non-native speakerlar, madaniy farqli talabalar) uchun sistematik adolatsizlikka olib keladi va bu ta'limdagi mavjud tengsizliklarni kuchaytirishi mumkin. O'zbekiston kontekstida ko'p tillilik va madaniy xilma-xillik tufayli global modellar mahalliy talabalar uchun mos kelmasligi ehtimoli yuqori, shuning uchun bias yanada keskinlashadi. Bias yumshatish choralari (ma'lumotlar diversifikatsiyasi, audit, fairness metrikalari, explainable AI, inson nazorati) samarali, lekin amalda qo'llash texnik, moliyaviy va institutsional to'siqlarga duch keladi, shuning uchun faqat texnik yechim yetarli emas — axloqiy me'yorlar va ta'lim siyosatini o'zgartirish zarur. Xulosa qilib aytganda, SI ta'lim baholashida muhim vosita bo'lib qoladi, ammo adolatli qo'llanilishi uchun mahalliy moslashtirish, shaffoflik va inson nazorati majburiy; keyingi tadqiqotlar empirik sinovlar va O'zbekiston uchun maxsus modellar ishlab chiqishga qaratilishi kerak.

Xulosa: Sun'iy intellekt ta'lim baholashida tezlik, izchillik va shaxsiylashtirish imkonini bersa-da, ishonchlilik va xolislik muammolari saqlanib qolmoqda: SI tizimlari inson bahosidan konsistentroq bo'lsa ham, algoritmik bias (nomutanosib ma'lumotlar, madaniy-lingvistik farqlar, "qora quti" effekti) tufayli ozchilik guruhlari uchun sistematik adolatsizlik yuzaga kelib, ta'limdagi tengsizlikni kuchaytirishi mumkin; O'zbekiston kontekstida ko'p tillilik va madaniy xilma-xillik sababli global modellar mahalliy talabalar uchun mos kelmasligi ehtimoli yuqori bo'lgani uchun mahalliy moslashtirish zarur; muammoni yumshatish uchun ma'lumotlar diversifikatsiyasi, muntazam audit, fairness metrikalari, shaffoflik (explainable AI) va inson nazorati majburiy choralar hisoblanadi; SI ta'limda katta salohiyatga ega bo'lsa-da, uning adolatli qo'llanilishi uchun texnik va institutsional yondashuvlar talab etiladi, kelajakda O'zbekiston uchun maxsus modellar ishlab chiqish va empirik sinovlar o'tkazish lozim.

Foydalanilgan adabiyotlar:

1. UNESCO. (2021). *AI and education: Guidance for policy-makers*. Paris: UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000376709>
2. U.S. Department of Education, Office of Educational Technology. (2023). *Artificial Intelligence and the Future of Teaching and Learning: Insights and Recommendations*. <https://tech.ed.gov/ai/ai-report/ai-report.pdf>
3. Yin, Z., et al. (2024). FairAIED: Navigating Fairness, Bias, and Ethics in Educational AI Applications. *arXiv preprint arXiv:2407.18745*. <https://arxiv.org/abs/2407.18745>
4. Idowu, J. A., et al. (2024). Investigating algorithmic bias in student progress monitoring. *Computers and Education: Artificial Intelligence*, 6, 100207. <https://www.sciencedirect.com/science/article/pii/S2666920X24000705>
5. Ferrara, E. (2023). Should ChatGPT be Biased? Challenges and Risks of Bias in Large Language Models. *arXiv preprint arXiv:2304.03738*. <https://arxiv.org/abs/2304.03738>
6. Bulut, O., & Cutumisu, M. (2024). Automated Essay Scoring and Fairness: Cultural and Linguistic Biases in Large Language Models. *Assessment in Education: Principles, Policy & Practice*. <https://www.tandfonline.com/doi/full/10.1080/0969594X.2024.2301234>