



AI ETIKASI: AQILLI TIZIMLAR QARORLARINI BOSHQARISH MASALALARI

Muallif: Fayziyeva Maftuna Ilxomovna¹

Affiliyatsiya: Iqtisodiyot va pedagogika universiteti, “Iqtisodiyot va muxandislik fanlari” kafedrası assistenti¹

DOI: <https://doi.org/10.5281/zenodo.17349718>

ANNOTATSIYA

Ushbu maqolada sun'iy intellekt (AI) texnologiyalarining jadal rivojlanishi natijasida yuzaga kelayotgan murakkab axloqiy, ijtimoiy va huquqiy muammolar keng ko'lamda tahlil qilinadi. AI tizimlarining qaror qabul qilish jarayonida shaffoflik, javobgarlik, adolat va inson nazorati bilan bog'liq masalalar chuqur yoritilgan. Shuningdek, maqolada aqlli tizimlarning qarorlarini boshqarish va tartibga solishning zamonaviy mexanizmlari, xalqaro tajribadan olinayotgan saboqlar hamda O'zbekiston uchun dolzarb bo'lgan yo'nalishlar ko'rib chiqiladi.

Kalit so'zlar: Sun'iy intellekt etikasi, shaffoflik, tarafkashlik, javobgarlik, inson nazorati.

KIRISH

Raqamli inqilob davri sun'iy intellekt texnologiyalarining keng qo'llanilishi bilan tavsiflanadi. AI sog'liqni saqlash, transport, ta'lim, moliya va davlat boshqaruvi kabi sohalarda katta o'zgarishlarni amalga oshirmoqda. Shu bilan birga, AI qarorlari inson hayotiga bevosita ta'sir qilishi tufayli etik va huquqiy masalalarni o'rganish dolzarbdır. Ushbu tadqiqotning maqsadi AI tizimlarining qarorlarini boshqarish mexanizmlari, ularning shaffoflik, javobgarlik va adolat tamoyillari asosida tahlil qilishdan iborat.

METODOLOGIYA

Tadqiqotda quyidagi metodlardan foydalanildi: - Tahliliy usul: xalqaro va milliy tajribalarni qiyosiy o'rganish; - Nazariy asos: AI etikasi bo'yicha ilmiy adabiyotlar, xalqaro standartlar va qonun hujjatlari; - Amaliy misollar: sog'liqni saqlash, moliya, sud tizimi va ijtimoiy tarmoqlardagi real holatlar tahlili. AI qarorlaridagi shaffoflik muammosi. Sun'iy intellekt tizimlarining qaror qabul qilish jarayonida eng muhim va bahsli masalalardan biri – bu **shaffoflik muammosidir**. AI ishlayotganida u qanday algoritmlar asosida yakuniy qarorga kelgani ko'pincha inson uchun tushunarsiz bo'lib qoladi. Shu sababli bu jarayon **“qora quti”** deb ataladi. Shaffoflik yetishmasligining sabablari

Birinchidan, zamonaviy AI modellarining asosiy qismi, ayniqsa chuqur o'rganish (deep learning) usuliga tayangan neyron tarmoqlar juda katta hajmdagi ma'lumotlarni qayta ishlaydi va millionlab parametrlar asosida hisob-kitob olib boradi. Bu parametrlarning har birini tahlil qilib, yakuniy qarorni qanday shakllantirganini oddiy foydalanuvchi tushunishi juda murakkab.

Ikkinchidan, AI modellarining aksariyatida **qaror qabul qilish jarayonini izohlab beruvchi maxsus modul** mavjud emas. Algoritmlar natija chiqaradi, ammo ushbu natijaning sababi haqida to'liq izoh berolmaydi.

Uchinchidan, AI texnologiyalari juda murakkab matematik va texnik asosga ega bo'lgani uchun mutaxassis bo'lmaganlar emas, ba'zan tajribali dasturchilar ham qarorlarni tushuntirib bera olmaydilar.

Shaffof bo'lmagan qarorlar jamiyatda bir qancha muammolarni keltirib chiqaradi:

1. **Ishonchning pasayishi** – Foydalanuvchilar va jamoatchilik AI natijalariga ishonmay qo'yishi mumkin.
2. **Xatolarni aniqlash qiyinligi** – Qaror noto'g'ri bo'lsa, nima uchun xato yuz berganini bilish mushkul.
3. **Javobgarlikni aniqlashdagi murakkablik** – Qaror uchun kim javobgar: algoritmnı ishlab chiqqan dasturchi, tizimdan foydalanuvchi tashkilot yoki platformaning o'zi?

Amaliy misollarda ko'rib chiqadigan bo'lsak. **Bank sektorida** kredit berish jarayonida AI tizimi ba'zi mijozlarning so'rovini qondiradi, boshqasini rad etadi. Agar qaror sabablari ko'rsatilmasa, bu adolatsizlik sifatida qabul qilinadi. **Tibbiyotda** AI rentgen tasviri asosida kasallikni aniqlaydi, ammo qaror ortidagi mantiq tushuntirilmasa, shifokor ushbu tavsiyani to'liq qabul qilmaydi. **Ijtimoiy tarmoqlarda** kontentni filtrlash yoki bloklash qarorlari ham ko'pincha tushuntirilmaydi, natijada foydalanuvchi o'z erkinliklari cheklangan deb hisoblaydi.

So'nggi yillarda shaffoflik muammosini bartaraf etish uchun bir qator yondashuvlar taklif qilinmoqda:

1. **Explainable AI (XAI)** – qarorlarni izohlaydigan, qaror qabul qilish jarayonini tushuntira oladigan algoritmlar ishlab chiqilmoqda.
2. **Huquqiy talablar** – Yevropa Ittifoqida foydalanuvchi o'ziga ta'sir qiluvchi AI qarori haqida tushuntirish olish huquqiga ega bo'lishini belgilovchi qonunlar ishlab chiqilgan.
3. **Ochiq manba (open-source) platformalar** – ayrim AI tizimlari kodini ochiq qilish orqali shaffoflikni oshirishga urinishlar mavjud.
4. **Etik standartlar** – xalqaro miqyosda AI tizimlari uchun axloqiy kodeks va standartlar ishlab chiqilmoqda.

Sun'iy intellekt tizimlari insonlar tomonidan yaratilgan ma'lumotlar asosida o'qitiladi. Shu sababli ularning qarorlari ko'pincha ushbu ma'lumotlar bazasida mavjud bo'lgan noxolislik (tarafkashlik) va xatolarni ham takrorlaydi. Bu hodisa **algoritmik tarafkashlik (algorithmic bias)** deb ataladi. Tarafkashlikning asosiy sabablari

1. **Ma'lumotlarning noto'g'ri yoki biryoqlama bo'lishi.** Agar o'quv ma'lumotlari to'liq bo'lmasa, tarixiy stereotiplar yoki noto'g'ri tasniflangan ma'lumotlar bo'lsa, AI ularni avtomatik ravishda takrorlaydi.
2. **Dizayn bosqichida insonning xatolari.** Modelni yaratish jarayonida ishlab chiquvchilar tomonidan noto'g'ri tanlangan parametrlar yoki algoritmlar ham tarafkashlikni keltirib chiqarishi mumkin.
3. **Texnologik imkoniyatlarning cheklanganligi.** Ba'zan ma'lumotlar hajmi yetarli emas yoki bir guruh ma'lumotlari boshqa guruhlarnikiga nisbatan ko'proq bo'ladi.

Bu modelni bir tomonlama qaror qabul qilishga olib keladi va buning oqibatida. AI qarorlari ayrim ijtimoiy, etnik yoki jinsiy guruhlarni kamsitishi mumkin. Noxolis qarorlar foydalanuvchilar orasida ishonchsizlikni kuchaytiradi.

NATIJALAR

Tadqiqot natijalariga ko'ra quyidagi xulosalar olindi:

AI texnologiyalari noto'g'ri ishlatilganda mavjud muammolarni yanada chuqurlashtiradi.

Misol uchun Sud qarorlarini baholash tizimlari (AQSh): Ayrim algoritmlar afro-amerikaliklarni yuqori xavfli deb baholab, oq tanlilarga nisbatan qattiqroq jazolar tavsiya qilgan.

Ishga qabul qilish algoritmlari: Ba'zi AI dasturlari erkaklar haqidagi tarixiy ma'lumotlar ko'pligi sababli erkak nomzodlarni afzal ko'rgan, ayollarni esa kamroq tanlagan.

Yuzni aniqlash texnologiyalari: Ko'pincha oq tanli yuzlarni to'g'ri tanib oladi, ammo qora tanli yoki osiyolik yuzlarda xato ehtimoli yuqori. Bunday kamchiliklarni kamaytirish uchun

1. **Ma'lumotlar bazasini to'g'ri shakllantirish.** O'quv ma'lumotlarining balansli, to'liq va sifatli bo'lishini ta'minlash.
2. **Audit va nazorat.** AI tizimlarining muntazam ravishda mustaqil ekspertlar tomonidan tekshirilishi.
3. **Algoritmik adolat (fairness) yondashuvlari.** Modelni yaratishda adolatni ta'minlaydigan matematik usullarni qo'llash.
4. **Xalqaro standartlar** Jahonda ishlab chiqilayotgan etik qoidalar va standartlarga amal qilish.

Javobgarlik va nazorat. Sun'iy intellektning keng joriy etilishi bilan bog'liq eng dolzarb savollardan biri — **AI qarorlariga kim javobgar bo'lishi va qarorlar ustidan qanday nazorat o'rnatilishi** masalasidir. AI tizimlari mustaqil qarorlar qabul qila boshlagan sari, ularning noto'g'ri yoki salbiy oqibatlarga olib kelgan qarorlarida javobgarlikni aniqlash qiyinlashib bormoqda. Asosiy savollar quyidagilardan iborat:

1. Xatoga **algoritmni yaratgan dasturchi** javob beradimi?
2. **Texnologiyani ishlatgan tashkilot** mas'ul bo'ladimi? yoki **foydalanuvchi** javobgar bo'ladimi?

Masalan, avtonom avtomobil avariyaiga sabab bo'lsa, kim javobgar: mashinani ishlab chiqargan kompaniya, dastur tuzgan muhandislar, yo'l harakati sharoitlari yoki avtomobil egasimi?

AI qarorlarini inson tomonidan nazorat qilish (human oversight) – xavfsizlikning asosiy kafolatidir. Nazoratning ikki asosiy shakli mavjud:

1. **Oldindan nazorat** – AI modellarini ishlab chiqish va o'qitish jarayonida sinov, audit, axloqiy ekspertiza o'tkazish.

Keyingi nazorat – AI qarorlarini kuzatib borish, ularni tushuntirish, kerak bo'lsa qarorlarni bekor qilish.

Xalqaro tajriba.

Sun'iy intellektning jadal rivojlanishi dunyo miqyosida yangi tartibga solish va axloqiy nazorat mexanizmlarini talab qilmoqda. Raqamli texnologiyalar tezkor joriy etilayotgani sababli ko'plab davlatlar AI etikasi bo'yicha maxsus strategiya va huquqiy hujjatlarni ishlab chiqmoqda. Quyida eng muhim yondashuvlar ko'rib chiqiladi.

Yevropa Ittifoqi sun'iy intellektni tartibga solish borasida eng faol hududlardan biridir. 2021-yilda ishlab chiqilgan **Artificial Intelligence Act** loyihasi quyidagilarni ko'zda tutadi:

- a) AI tizimlarini xavf darajasiga ko'ra toifalarga ajratish;
- b) Yuqori xavfli AI modellarida inson nazoratini majburiy qilish;

- c) Qaror qabul qilish jarayonida shaffoflikni ta'minlash;
- d) Foydalanuvchilar huquqlarini himoya qilish.

Bundan tashqari, Yevropa Komissiyasi tomonidan **“Etik tamoyillar asosida ishonchli AI”** (Ethics Guidelines for Trustworthy AI) hujjati qabul qilingan.

AQShda AI texnologiyalarini to'liq federal qonun bilan emas, balki soha bo'yicha **tavsiya va standartlar orqali** tartibga solish amaliyoti mavjud. 2023-yilda NIST (National Institute of Standards and Technology) **AI Risk Management Framework**ni taqdim etdi. Ushbu hujjat:

- AI xavflarini boshqarish;
- Algoritmardagi noxolislikni kamaytirish;
- AI qarorlarini izohlash va nazorat mexanizmlarini takomillashtirishga qaratilgan.

Kanadada hukumat davlat xizmatlarida AI texnologiyalarini joriy etishdan oldin **Algorithmic Impact Assessment** (algoritmik ta'sir bahosi) usulini qo'llaydi. Bu usul:

- AI loyihasining jamiyatga ta'sirini baholash;
- Shaffoflik va javobgarlikni oshirishni maqsad qiladi.

Yaponiya **“Jamiyat 5.0”** konsepsiyasi doirasida inson markazida sun'iy intellektni rivojlantirish siyosatini yuritmoqda. Bu yondashuvda:

- Texnologiya inson hayotini qulaylashtirishga xizmat qilishi;
- Avtomatlashtirish jarayonida inson qadriyatlari birinchi o'rinda bo'lishi ko'zda tutiladi.

2021-yilda UNESCO tomonidan **“AI etikasi bo'yicha xalqaro tavsiyalar”** qabul qilindi. Ushbu hujjat:

- AI texnologiyalaridan foydalanganda inson huquqlarini ta'minlash;
 - Adolat, shaffoflik va javobgarlik tamoyillarini mustahkamlash;
 - Xalqaro hamkorlikni kengaytirishga qaratilgan.
- Dunyo tajribasi shuni ko'rsatadiki:

AI rivoji bilan birga **inson huquqlari va etik masalalar birinchi o'ringa chiqmoqda.**

Har bir davlat o'z sharoitidan kelib chiqib tartibga solish yondashuvini tanlamoqda.

Global standartlar va tamoyillar asta-sekin shakllanib bormoqda, bu esa barcha mamlakatlar uchun yagona etik maydon yaratishga yordam beradi.

O'zbekiston tajribasi. O'zbekiston so'nggi yillarda raqamli texnologiyalar, shu jumladan sun'iy intellekt (AI) sohasini rivojlantirishga katta e'tibor qaratmoqda. Davlat siyosatining ustuvor yo'nalishlaridan biri – **raqamli iqtisodiyotni** shakllantirish va barcha sohalarga raqamli texnologiyalarni keng joriy etishdir. Shu bilan birga, AI bilan bog'liq huquqiy, axloqiy va xavfsizlik masalalari ham asta-sekin kun tartibiga chiqmoqda.

2020-yilda qabul qilingan **“Raqamli O'zbekiston – 2030”** davlat dasturi:

Raqamli texnologiyalarni barcha sohalarga keng joriy etish;

Axborot xavfsizligi va shaxsiy ma'lumotlarni himoya qilish mexanizmlarini mustahkamlash;

Sun'iy intellektni rivojlantirishga doir ilmiy markazlar tashkil etishni maqsad qilgan.

Bu dastur AI sohasidagi tadqiqotlar va innovatsion loyihalar uchun umumiy strategik asos yaratmoqda.

So'nggi yillarda qator normativ-huquqiy hujjatlar qabul qilindi:

- “Shaxsga doir ma'lumotlar to'g'risida”gi Qonun (2020-yilda yangilangan);
- Axborot xavfsizligi konsepsiyasi;
- Raqamli iqtisodiyot va kiberxavfsizlikni rivojlantirish strategiyalari.

Mazkur hujjatlar bevosita AI faoliyatini to'liq tartibga solmasa-da, **shaxsiy ma'lumotlar bilan ishlashda himoya choralarini** belgilab beradi va kelajakda AI texnologiyalari uchun normativ asos bo'lib xizmat qiladi.

- ✓ Toshkent shahrida **Sun'iy intellekt markazi** faoliyat yuritmoqda.
- ✓ Turli universitetlarda AI va ma'lumotlar tahlili bo'yicha yangi ta'lim yo'nalishlari ochildi.

Bank, tibbiyot va davlat xizmatlarida AI asosidagi dasturiy yechimlar joriy qilinmoqda. Hozirgi bosqichda O'zbekistonda AI texnologiyalarining keng qo'llanilishi bilan bog'liq bir qator masalalar dolzarbligicha qolmoqda.

MUHOKAMA

Natijalar shuni ko'rsatadiki, AI etikasi nafaqat texnik, balki huquqiy, falsafiy va ijtimoiy masalalar bilan bevosita bog'liqdir. Inson nazorati AI qarorlarining xavfsizligini ta'minlashda muhim ahamiyat kasb etadi. Xalqaro tajribadan foydalanib, O'zbekiston sharoitiga mos qonunchilik va etik standartlar ishlab chiqilishi zarur. Kelajakda AI mustaqil qaror qabul qiluvchi subyektga aylanishi ehtimoli mavjud bo'lgani sababli, bu borada inson qadriyatlari ustuvor bo'lishi lozim.

XULOSA

Kelgusida aqlli tizimlar nafaqat yordamchi, balki mustaqil qaror qabul qiluvchi subyekt sifatida faoliyat yuritishi ehtimoli katta. Shunday vaziyatda AI etikasi nafaqat texnik fan, balki huquqshunoslik, sotsiologiya va falsafa bilan ham uzviy bog'liq holda rivojlanadi. Bu esa insonning texnologiyalar ustidan yakuniy nazoratini saqlab qolishning yagona yo'li bo'lib qoladi. Xulosa qilib aytganda, AI texnologiyalari insoniyat taraqqiyotining yangi bosqichini boshlab berdi. Ammo ularning qarorlarini boshqarish va nazorat qilish, inson huquqlarini himoya qilish hamda adolatni ta'minlash jamiyat oldida eng muhim vazifa sifatida turibdi. Kelgusida texnologiya imkoniyatlaridan oqilona foydalanish orqali AI jamiyatning barqaror rivojlanishiga xizmat qilishi mumkin.

FOYDALANILGAN ADABIYOTLAR RO'YXATI

1. Floridi, L., & Cows, J. (2021). A Unified Framework of Five Principles for AI in Society. Harvard Data Science Review.
2. Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. Nature Machine Intelligence, 1(9), 389–399.
3. European Commission. (2021). Proposal for a Regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).
4. Bostrom, N. (2014). Superintelligence: Paths, Dangers, Strategies. Oxford University Press.
5. Bryson, J. J. (2019). The Past Decade and Future of AI's Impact on Society. In Towards a New Enlightenment? A Transcendent Decade.
6. Ilkhomovna F. M. (2025). Pedagogical Conditions of Realisation of Educational Functions in The Process of Educational Activity in Universities. Academia Open, 10(1).