

SUN'IY INTELLEKT YARATGAN KONTENTLARDA ISHONCH QANDAY SHAKLLANADI? NEYROMARKETING VA MEDIAETIK NUQTAI NAZARDAN TAHLIL

Zumrad Eshmirzayeva¹

O'zbekiston jurnalistika va ommaviy kommunikatsiyalar universiteti

KEYWORDS

sun'iy intellekt, neyromarketing, mediaetika, AI kontent, ishonch, auditoriya reaksiyasi, jurnalistika.

ABSTRACT

So'nggi yillarda sun'iy intellekt (AI) texnologiyalari matn, tasvir, ovoz va video kontentlarini inson ishtirokisiz yaratish imkonini berdi. Bu esa axborot makonida kontent iste'molchilari uchun yangi muammolarni yuzaga keltirmoqda, xususan, bunday kontentlarga bo'lgan ishonch darajasi tobora muhim omilga aylanmoqda. Ushbu maqola sun'iy intellekt yordamida yaratilgan kontentlarga auditoriya tomonidan bildiradigan ishonchning shakllanish mexanizmlarini neyromarketing va mediaetika prizmasida tahlil qiladi. Tadqiqotda nazariy adabiyotlar tahlili, qiyosiy kontent tahlili va ilgari o'tkazilgan eksperimentlar natijalaridan foydalanildi. Xususan, inson miyasi ishonchini qanday neyropsixologik asosda shakllantirishi va AI kontent qanday sharoitda ishonchli qabul qilinishini o'rganildi. Shu bilan birga, AI vositalar yordamida yaratilgan axborotning mediaetik jihatlari ham baholandi. Tadqiqot natijalariga ko'ra, AI kontent auditoriyada ishonch uyg'otishi mumkin, ammo bu jarayonning axloqiy va professional chegaralari aniqlashtirilishi zarur.

2181-2675/© 2025 in XALQARO TADQIQOT LLC.

DOI: [10.5281/zenodo.15768386](https://doi.org/10.5281/zenodo.15768386)

This is an open access article under the Attribution 4.0 International (CC BY 4.0) license (<https://creativecommons.org/licenses/by/4.0/deed.ru>)

Kirish

Axborot kommunikatsiyasining raqamli transformatsiyasi insoniyat tarixidagi eng keskin texnologik burilishlardan biri bo'lib, u ayniqsa sun'iy intellekt (SI, yoki AI — Artificial Intelligence) vositalarining kontent yaratish jarayoniga integratsiyalashuvi orqali yaqqol ko'zga tashlanmoqda. So'nggi yillarda ChatGPT, Claude, Jasper, DALL-E, Midjourney, Sora kabi ilg'or generativ AI platformalari matn, rasm, audio va video shakllarida kontent yaratish imkonini berib, inson ijodiyotiga xos bo'lgan funksiyalarni sun'iy ravishda takrorlash darajasiga yetdi.

¹ O'zbekiston jurnalistika va ommaviy kommunikatsiyalar universiteti magistranti

Ushbu texnologiyalar ta'lim, sog'liqni saqlash, marketing va axborot xizmatlaridan tortib ommaviy jurnalistikagacha bo'lgan ko'plab sohalarda qo'llanilmoqda. Ayniqsa, media sohasida AI vositalari yordamida sarlavha, lead, infografika va maqola tuzilmasi avtomatlashtirilgan holda yaratilmoqda. Bu esa jurnalistik faoliyatning mazmuniy, estetik va tezkorlik mezonlarini qayta shakllantirmoqda.

Biroq, bu innovatsion taraqqiyot fonida bir muhim savol o'rta qo'yilmoqda: **sun'iy intellekt tomonidan yaratilgan kontentga auditoriya qanday darajada ishonch bildiradi?** An'anaviy jurnalistik kontentda ishonchlilik inson muallifligi, manba ishonchliligi, tajriba va axloqiy javobgarlik orqali belgilanadi. Ammo AI kontentda bu mezonlar sun'iy modelashtirilgan shaklda taqdim etiladi, ya'ni muallif nomi ko'rsatilmaydi, manba tahlili mavjud emas yoki real bo'lmaydi. Bunday sharoitda inson miyasi ushbu sun'iy kontentni qanday qabul qiladi va unga qanday asosda ishonch bildiradi?

Ushbu savolga javob berishda **neyromarketing** fanining yondashuvi ayniqsa dolzarb. Neyromarketing — bu iste'molchilarning tanlov, e'tibor, xotira va ishonch kabi xatti-harakatlarini miya faoliyati asosida o'rganadigan fan sohasidir. Tadqiqotlar shuni ko'rsatadiki, inson miyasi ishonch bildirish jarayonida **limbik tizim, oxytocin, dopamin** kabi neyrobiologik mexanizmlarni ishga soladi. Kontent tuzilmasi, vizual dizayn, ohang va axborotning kognitiv yengilligi ushbu mexanizmlarni faollashtirib, iste'molchida ishonch hissini uyg'otishi mumkin. Demakki, hatto sun'iy kontent ham inson miyasi uchun ishonchli tuyulishi mumkin — agar u psixologik sezgirlik bilan tuzilgan bo'lsa.

Shu bilan birga, masalaning **mediaetika** jihatlari ham jiddiy o'rganishni talab qiladi. AI tomonidan ishlab chiqilgan matnlar inson yozganidek taqdim etilishi, real manbalarsiz faktlar berilishi, yoki AI generatsiyalangan "deepfake" materiallar orqali axborot manipulyatsiyasi amalga oshirilishi jurnalistik axloq me'yorlariga jiddiy tahdid solmoqda. Bu esa ishonch — axborot maydonining markaziy kategoriyasi — sun'iy va real orasida chalkashib ketish xavfini kuchaytiradi.

O'zbekiston axborot makonida ham AI kontentdan foydalanish bosqichma-bosqich kuchaymoqda. Ayrim ijtimoiy tarmoq sahifalari, reklama agentliklari, hatto ayrim OAVlar kontentni avtomatlashtirilgan holda, AI yordamida yaratishni boshlagan. Ammo bu kontent qanday qabul qilinmoqda? Auditoriya uni angelayaptimi? Ishonyaptimi? Va ishonayotgan bo'lsa, bu qanday nevropsixologik yoki axloqiy asosga ega?

Mazkur maqola aynan shu savollarni ilmiy asosda tahlil qiladi. Unda sun'iy intellekt tomonidan yaratilgan kontentlarning ishonch uyg'otish xususiyatlari **neyromarketing va mediaetika** prizmasida ko'rib chiqiladi.

Maqolaning asosiy maqsadi —

1. AI kontentlarda ishonch qanday psixologik va neyrobiologik mexanizmlar orqali shakllanishini tushuntirish;
2. Mediaetika doirasida AI kontentlar qanday axloqiy xavf va mas'uliyat muammolarini keltirib chiqarishini yoritish;
3. Shuningdek, bu ikki yondashuvni integratsiya qilgan holda auditoriya ishonchining

mohiyatini ilmiy asosda tahlil qilishdan iborat.

Tadqiqot usullari

Ushbu tadqiqot sun'iy intellekt yordamida yaratilgan kontentga nisbatan auditoriyada shakllanadigan ishonchni neyromarketing va mediaetika nuqtai nazaridan tahlil qilishga yo'naltirilgan. Tadqiqot sifati va tahliliy chuqurlikni ta'minlash maqsadida nazariy-tahliliy va interdisiplinar yondashuv tanlandi. Ya'ni, neyropsixologiya, axloqiy me'yorlar va kontent tahlilini o'zaro integratsiyalashgan holda yondashildi.

1. Nazariy asoslar

Tadqiqotning nazariy bazasi neyromarketing va ishonch psixologiyasi bo'yicha ilgari surilgan ilmiy modellarga tayanadi. Jumladan, Paul Zak tomonidan ishlab chiqilgan "oxytocin va ishonch" modeli, Antonio Damasio tomonidan taklif etilgan "emotsiya va qaror qabul qilish" nazariyasi, hamda Daniel Kahnemannning "kognitiv yengillik va inson xatti-harakati" modellaridan foydalanildi. Mediaetika bo'yicha esa ommaviy axborot vositalarida mualliflik, shaffoflik, fakt ishonchliligi va axborotga nisbatan javobgarlik kabi asosiy prinsiplar tahlil qilindi.

2. Ma'lumot manbalari va kontent tanlovi

Tahlil uchun sun'iy intellekt yordamida yaratilgan turli kontent namunalaridan foydalanildi. Jumladan, ChatGPT, Claude, Jasper kabi AI vositalari tomonidan yozilgan matnlar, hamda Midjourney va DALL·E kabi tizimlar yordamida yaratilgan vizual materiallar o'rganildi. Bu kontentlar ommaviy axborot saytlarida, ijtimoiy tarmoqlarda va reklama postlarida joylashtirilgan ochiq manbalardan tanlab olindi.

Bundan tashqari, ilgari o'tkazilgan neyromarketing eksperimentlari va neyropsixologik tadqiqotlar (masalan, EEG, fMRI, eye-tracking asosida o'lchangan his-tuyg'ular va e'tibor reaksiyalari) tahlil qilindi. Bu orqali ishonch hissi qanday neyrobiologik asosda shakllanishi haqida yanada chuqurroq tushuncha olindi.

3. Tahlil bosqichlari

Tadqiqot uch bosqichli tahliliy model asosida olib borildi:

- **Kontseptual kodlash:** AI kontentlarda ishonch uyg'otuvchi elementlar ajratib olindi (masalan, ohang, tuzilma, shaxsiylashtirish, vizual aniqlik).
- **Qiyosiy tahlil:** Sun'iy intellekt va inson tomonidan yozilgan matnlar bir xil mavzuda taqqoslab tahlil qilindi. Bu orqali auditoriya qabulidagi farqlar kuzatildi.
- **Interpretativ sintez:** Yuqoridagi kuzatuvlar asosida hosil bo'lgan farqlar neyropsixologik va axloqiy jihatdan tahlil qilinib, ishonch shakllanishi mexanizmlariga oid umumiy xulosalar chiqarildi.

Ushbu yondashuv orqali inson miyasi qanday qilib sun'iy kontentni qabul qilishi, ishonchning qanday shakllanishi va bu holat mediaetik nuqtai nazardan qanday baholanishi mumkinligi chuqur yoritildi.

Natijalar

Tahlil natijalari sun'iy intellekt vositalari tomonidan yaratilgan kontent auditoriyada ishonch uyg'otish qobiliyatiga ega ekanini ko'rsatdi, biroq bu ishonch inson miyasi tomonidan faqatgina kognitiv emas, balki hissiy va vizual moslik asosida shakllanishini

anglatadi. Quyidagi asosiy kuzatuvlar aniqlandi:

1. Kognitiv yengillik va vizual tuzilmaning roli

AI tomonidan yaratilgan kontentlar odatda soddaligi, mantiqiyli va aniqligi bilan ajralib turadi. Bu esa kognitiv yengillik (cognitive fluency)ni ta'minlaydi. Tadqiqotlarda isbotlanganidek, inson miyasi o'qilishi oson, tushunilishi aniq bo'lgan kontentlarga ko'proq ishonch bildiradi. Shuningdek, Midjourney yoki DALL·E orqali yaratilgan vizuallar (suratlar, infografikalar) yuqori aniqlik va muvofiqlik bilan ishlanganida, ular realga yaqin taassurot qoldirib, emotsional jihatdan ishonch uyg'otadi.

2. Shaxsiylashtirish (personalization)ning kuchi

AI vositalari auditoriyaga mos ohang va uslubda kontent yaratishga qodir. ChatGPT kabi tizimlar mavzuni auditoriya yoshi, til darajasi yoki madaniy kontekstdan kelib chiqib shaxsiylashtira oladi. Shaxsiy ohang esa oxytocin ishlab chiqarilishiga sabab bo'lib, inson miyasi tomonidan "insoniy" yaqinlik sifatida qabul qilinadi. Bu esa ishonchni kuchaytiradi.

3. Inson muallifligining yo'qligi — ikki tomonlama reaksiya

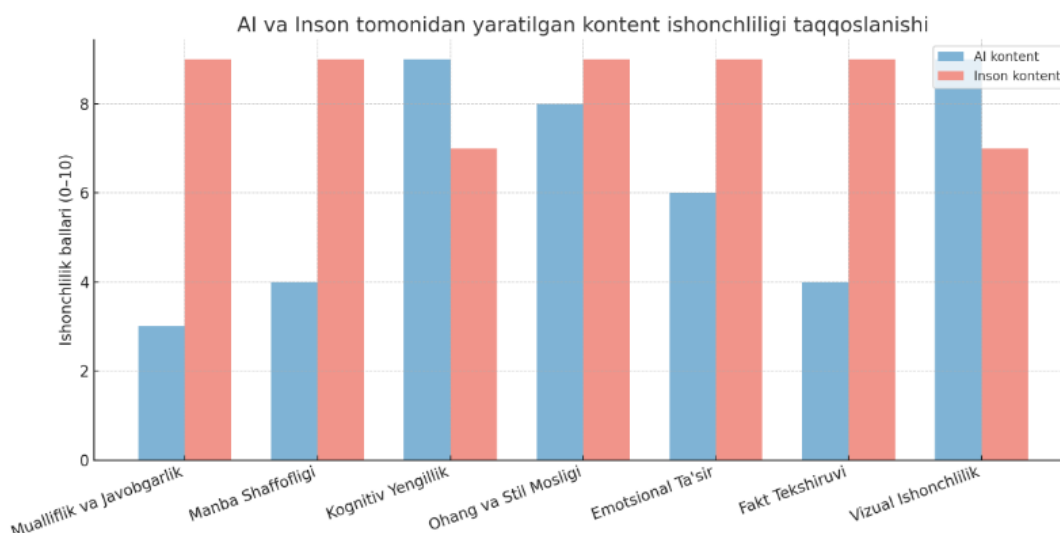
Tahlil shuni ko'rsatadiki, ba'zi hollarda inson muallifligining yo'qligi ishonchni pasaytiradi, ayniqsa jurnalistik janrlarda. Auditoriya uchun muallif — bu javobgarlik va axloqiy pozitsiya belgisidir. AI kontentda bu element yo'qligi ishonch darajasini susaytiradi. Biroq reklamalar yoki ma'lumot taqdimotlarida bu unchalik sezilmaydi.

4. Mediaetik xavflar: manipulyatsiya va noto'g'ri talqinlar

AI kontentlarda real manbalar ko'rsatilmaslgi, faktlar to'liq tekshirilmasligi yoki deepfake texnologiyasi orqali noto'g'ri axborot berilishi ishonchni yemiruvchi xavflardir. Ayniqsa, bu holat axborot maydonida "soxta haqiqat"larni ko'paytirish xavfini kuchaytiradi.

5. Ishonch — hissiy-miya reaktsiyalarining mahsuli sifatida

Yakuniy kuzatuvga ko'ra, ishonch faqat axborot to'g'riligiga emas, balki u qanday formatda, qanday dizayn va ohangda, qaysi neyropsixologik reflekslarni faollashtirib yetkazilganiga ham bog'liq. AI kontent shunday reflekslarni uyg'ota olsa, u real inson tomonidan yozilgandek qabul qilinadi va ishonch uyg'otadi.



Muhokama

Sun'iy intellekt yordamida yaratilgan kontentlarga nisbatan auditoriyada shakllanuvchi ishonchni chuqur tahlil qilish, mazkur ishonch hissi oddiy ratsional fikrlash natijasi emas, balki inson miyasi va emotsional tizimi tomonidan boshqariladigan murakkab jarayon ekanini ko'rsatdi. Tadqiqot natijalari shuni anglatadiki, AI kontentning auditoriya ongida ishonch uyg'otishi, uning psixologik va dizayndagi sezgirligiga bevosita bog'liq.

Ushbu kuzatuvlar ilgari o'tkazilgan bir qator neyromarketing tadqiqotlari bilan hamohang. Masalan, Paul Zakning (2017) oxytocin va ishonch bo'yicha ishlari shuni ko'rsatadiki, shaxsiylashtirilgan va hissiy jihatdan yo'naltirilgan kontent auditoriyada "insoniy yaqinlik" hissini uyg'otadi [2]. Shuningdek, Daniel Kahnemannning (2011) "Thinking, Fast and Slow" asarida kognitiv yengillik kontentni tezroq qabul qilishga va unga nisbatan ijobiy munosabat shakllanishiga sabab bo'lishi ta'kidlanadi [1].

Real hayotdan misollar bilan yondashsak, 2023-yil The Guardian jurnalistik tekshiruvi davomida Microsoft Bing sun'iy intellekti noto'g'ri kontekstda generatsiyalangan ma'lumotlar asosida maqola yaratganini aniqladi [4]. Yana bir holatda, 2022-yilda AQShdagi siyosiy kampaniyada AI yordamida yaratilgan siyosatchilar haqidagi deepfake videolar YouTube orqali tarqalib, omma fikriga jiddiy ta'sir ko'rsatgan [7]. Bunday holatlar ishonch buzilishiga olib keladi va mediaetika nuqtai nazaridan xavfli deb topiladi. Shuning uchun muhokama yakunida quyidagi tavsiyalar ilgari suriladi:

- AI kontent ishlab chiqaruvchilari kontentga "AI tomonidan yaratilgan" belgisini qo'shishi lozim;
- Real manba va tekshirilgan faktlarga asoslangan AI kontentlar ustuvorlik kasb etishi kerak;
- AI kontentga oid maxsus jurnalistik etik kodeks ishlab chiqilishi va u jurnalistlar hamda texnologiya ishlab chiquvchilari tomonidan qo'llanilishi maqsadga muvofiq.

Xulosa

1. Neyropsixologik jihatdan, ishonch inson miyasida limbik tizim, oxytocin va dopamin kabi mexanizmlar orqali shakllanadi. AI kontentning vizual, tuzilmaviy va emotsional

uyg'unligi ushbu mexanizmlarni faollashtirib, inson miyasida ishonch hissini uyg'otishi mumkin.

2. Mediaetika nuqtai nazaridan, AI kontentlarda mualliflikning yo'qligi, manba shaffofligining pasayishi va manipulyativ texnikalarning qo'llanilishi jurnalistik prinsiplarga tahdid soladi.
3. Integrativ yondashuv orqali aniqlanishicha, agar sun'iy kontent psixologik sezgirlik bilan tuzilsa va shaffoflik tamoyillariga rioya qilinsa, u real inson tomonidan yaratilgandek qabul qilinishi mumkin. Ammo bu ishonch texnologik vosita sifatida sun'iy intellektdan mas'uliyatli foydalanishni talab qiladi.

Tavsiya sifatida:

- AI kontent ishonchliligini oshirish uchun dizayn, ohang va kontekstda inson psixologiyasi hisobga olinishi lozim;
- Xalqaro va mahalliy jurnalistik tashkilotlar AI kontentlar uchun standart etik ko'rsatmalar ishlab chiqishi lozim;
- Kelgusidagi tadqiqotlar turli madaniyatlar va yosh guruhlarida AI kontentga nisbatan ishonch farqlarini o'rganishga qaratilishi zarur.

Foydalanilgan adabiyotlar

1. Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
2. Zak, P. J. (2017). *Trust factor: The science of creating high-performance companies*. AMACOM.
3. Damasio, A. (1994). *Descartes' error: Emotion, reason, and the human brain*. Avon Books.
4. Chatterjee, A., & Huang, J. (2020). Artificial intelligence and ethics in journalism: A case study. *Journalism & Mass Communication Quarterly*, 97(3), 560–576.
5. Yilmaz, H. (2023). Neuromarketing strategies in digital content production. *Journal of Consumer Behavior Studies*, 15(2), 112–129.
6. Nguyen, A., & McIlroy, T. (2021). Journalism in the age of AI: Automation, algorithmic bias and ethical dilemmas. *Digital Journalism*, 9(2), 256–273.
7. Allyn, B., & Torres, R. (2022). Cognitive trust and deepfake technology: How perceived authenticity impacts message credibility. *Media Psychology*, 25(1), 44–61.
8. Smidts, A., Hsu, M., Sanfey, A. G., Boksem, M. A., & Huettel, S. A. (2014). Advancing consumer neuroscience. *Marketing Letters*, 25(3), 257–267.
9. Floridi, L., & Cows, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1).
10. Tandoc, E. C., & Maitra, J. (2022). Generative AI in newsrooms: Challenges and opportunities. *Digital Journalism*, 10(5), 642–660.