



SUN'IY INTELLEKT TIZIMLARIDA ALGORITMIK TAFSIVNING FOYDALANUVCHI ONGIGA TA'SIRI: KOGNITIV TARAFKASHLIK VA AXBOROT FILTRI FENOMENI

Muallif: Shodiboyev Shohruh Shuhrat o'g'li¹

Affiliyatsiya: Xalqaro Nordik universiteti o'qituvchisi¹

DOI: <https://doi.org/10.5281/zenodo.17355800>

ANNOTATSIYA

Ushbu maqolada algoritmik tavsiyalar va kognitiv tarafkashlik o'rtasidagi o'zaro bog'liqlik tahlil qilinadi. Sun'iy intellekt asosidagi algoritmik tizimlar foydalanuvchilarga axborotni ularning qiziqishlari va ilgari qabul qilingan qarashlariga mos holda taqdim etadi. Bu jarayon qulaylik yaratish bilan birga, inson tafakkurida tasdiqlovchi tarafkashlik (confirmation bias), birinchi axborotga yopishish (anchoring effect) va ramka effekti (framing effect) kabi kognitiv og'ishlarni kuchaytiradi. Tadqiqot Nordik universitetining 1-3 kurs talabalari o'rtasida o'tkazilib, 100 nafar respondent ishtirok etdi. Eksperimental usulda turli algoritmik axborot oqimlari bilan tanishtirilgan uch guruh talabalarning qaror qabul qilishdagi tafovutlari o'rganildi. Cognitive Reflection Test (CRT), Framing Bias Scale va Media Literacy Scale metodikalari yordamida o'lchovlar amalga oshirildi. Natijalar shuni ko'rsatdiki, algoritmik filtrlash talabalarda tanlov erkinligini kamaytirib, mavjud e'tiqodlarni mustahkamlovchi kognitiv og'ishlarni kuchaytiradi. Shuningdek, media savodxonlik darajasi yuqori bo'lgan talabalar algoritmik tarafkashlik ta'siriga kamroq berilishi aniqlangan. Maqola algoritmik tizimlarning psixologik oqibatlarini tahlil qilib, kognitiv og'ishlarni kamaytirish bo'yicha tavsiyalarni ilgari suradi.

Kalit so'zlar: sun'iy intellekt tizimlari, algoritmik tafsiv, kognitiv tarafkashlik, axborot filtri, axborot pufagi, foydalanuvchi ongi, media algoritmlari, raqamli muhit, e'tibor iqtisodiyoti, psixologik ta'sir.

KIRISH

So'nggi yillarda internet va ijtimoiy tarmoqlardagi axborot oqimi sezilarli darajada algoritmlarga tayanmoqda. Foydalanuvchi qiziqishlari va xatti-harakatlariga moslashtirilgan algoritmik filtrlar (tafsiv) axborotni shaxsiylashtirib taqdim etadi. Natijada har birimiz axborot pufagi (filter bubble) deb ataladigan holatga tushib qolish xavfi mavjud. Ya'ni, izohlaylik: shaxsiy qidiruv tarixi va xatti-harakatlarimiz asosida tuzilgan algoritmlar bizga bir xil e'tiqod va qarashga mos keluvchi ma'lumotlarni ko'proq ko'rsatib, qarama-qarshi yoki turlicha qarashlarni yashirishi mumkin. Eli Pariser (2011) bu holatni keltirib chiqargan va "axborot pufagi" atamasini ilgari surgan; u fikricha, algoritmik filtrlar bir tomonlama moslashtirishga olib kelib, odamlarni intellektual jihatdan izolyatsiya qiladi va ijtimoiy bo'linishga sabab bo'lishi mumkin [1, 23b]. Shunday qilib, foydalanuvchining dunyoni qabul qilish doirasi torayadi – ko'proq tanish va qulay kontent ko'rib, boshqacha fikrlar qoplangan bo'ladi.

Bu jarayon shunchaki texnologik yangilikdan tashqari, inson psixologiyasiga ham ta'sir qiladi. Algoritmik tavsiyalar mavjud ilgari e'tiqodlarimizni mustahkamlab, kognitiv tarafkashliklarni kuchaytirishi mumkin. Kognitiv tarafkashlik (ya'ni oldindan mavjud qarashni tasdiqlovchi ma'lumotlarni afzal ko'rish) ongsiz ravishda yuzaga kelib, inson qarorlarini toraytiradi. Tadqiqotlar shuni ko'rsatmoqdaki, onlayn muhitda tasdiqlovchi tarafkashlik (confirmation bias) kuchli ifoda topadi: odamlar o'z oldidagi qarashlarga mos keluvchi ma'lumotlarni ko'proq qabul qiladi va qarama-qarshi fikrlarni mensimaslikka moyil bo'ladi. Bunday tarafkashlik onlayn platformalarda emotsional va qo'zg'aluvchi kontent bilan birga algoritm tomonidan ham oshib boradi: algoritmik personalizatsiya odamlarni o'zlariga o'xshash fikrlar bilan «echo chamber» va axborot pufagiga mahkam o'rab qo'yadi [2].

ASOSIY QISM

Algoritmik tafsiv ijtimoiy media va qidiruv tizimlarida insonlarning harakatlariga binoan axborotni tanlash va saralash tizimidir. Masalan, Google yoki YouTube algoritmlari foydalanuvchi qiziqishlarini aniqlab, unga mos keluvchi ma'lumotlarni birinchi o'ringa qo'yadi. Bu jarayon foydalanuvchi uchun qulay bo'lsa-da, negizida «ma'lumot pufagi» paydo bo'lishiga olib kelishi mumkin. Pariser fikriga ko'ra, foydalanuvchining shaxsiy xulq-atvori va tarixini o'rganib, algoritmlar shaxsiylashtirilgan kontentni taklif qiladi va bu jarayon uni boshqacha axborotlardan uzib qo'yadi.

Bundan tashqari, zamonaviy tadqiqotlar «informatcion kokonglar (information cocoons)» atamasini ham qo'llaydi. Masalan, bir tadqiqotda ta'kidlanishicha, algoritmik filtrlash orqali axborotga kirish kengayib borayotgan bo'lsada, foydalanuvchilar shu bilan birga o'zlarini nazariy jihatdan erkin his etmay, aslida «informatcion kokong» ichida qolish xavfi mavjudligi aytiladi. Ya'ni, axborotni algoritmlar orqali filtrlash odamlarning dunyoqarashini kengaytirmoqchi ekan, amalda ularni bitta ma'lumot oqimiga qamash va yangicha qarashlardan ajratish xavfi bor. Natijada, algoritm tizimining «neytral» ekanligi biz o'ylagandek bo'lmasligi, balki axborot ierarxiyasini shakllantirishga xizmat qilishi ehtimoli yuqori.

Kognitiv tarafkashliklar – inson ongiga xos bo'lgan, qarorlarimizni va tanlovlarimizni toraytiruvchi sistematik xatoliklar yoki moyilliklardir. Ushbu maqolada ayniqsa quyidagi uch main tarafkashlikka e'tibor qaratamiz:

- Tasdiqlovchi tarafkashlik (Confirmation Bias): odamlar o'zining mavjud e'tiqodlarini tasdiqlovchi ma'lumotlarni izlash va ularga ishonishga moyil bo'lib, qarama-qarshi ma'lumotni mensimaslik tendensiyasini anglatadi.[3] Ya'ni, ilgari qabul qilingan fikrlarni tasdiqlovchi yangi dalillar topishga harakat qiladi va o'z fikriga zid bo'lgan ma'lumotni idrok etishda xato qiladi.
- Birinchi ko'rgan natijaga yopishish (Anchoring Effect): qaror qabul qilishda yoki baho berishda ilk ko'rsatilgan raqam yoki ma'lumot (anchor) ta'siri ostida qolish holatidir. Masalan, narx bo'yicha muzokaralarda dastlab aniq yoki yuqori raqam berilsa, keyingi baholar bunga nisbatan qabul qilinadi. Ushbu effekt shamollikki, inson miyasida dastlab tushgan informatsiya me'yor sifatida qabul qilinadi va keyingi qarorlar shunga bo'ysunadi.
- Ramka effekti (Framing Effect): tanlov variantlari va qarorlar taqdim etilish usuliga ko'ra odamlar turlicha qaror qabul qilishidir. Masalan, ikki xil holatda ham natijasi bir xil bo'lgan vaziyatlardan biri «yaxshi tomonlar» bilan, ikkinchisi «yomon tomonlar» bilan taqdim etilganida, odamlarning qarori farqlanishi

mumkin. Shunday qilib, bir xil axborot ijobiy yoki salbiy tusda berilganda, odamlarning unga munosabatida farqlik yuz beradi [4].

Yuqoridagi tarafkashliklar ongli va ongsiz ravishda har qanday ma'lumot qabul qilish jarayonida namoyon bo'lishi mumkin. Ularning majmuaviy ta'siri esa, ayniqsa onlayn-muhitda, algoritmik tavsiyalar bilan birlashganda yanada kuchayadi. Tadqiqotlar ko'rsatadiki, algoritmlar odamlarning allaqachon mavjud bo'lgan tarafkashliklarini mustahkamlash yo'lida ishlaydi. Misol uchun, odamlar ilgari e'tiqod qilishgan ma'lumot turidan chetga chiqmaydi va xuddi shu turdagi yangi kontentni ko'proq qabul qiladi. Onlayn muhitda esa bu tendensiya kuchayadi: chunki algoritmlar foydalanuvchining qiziqishini uyg'otadigan kontentni berishga intiladi, ko'pincha bu esa mavjud tarafkashlikni mustahkamlaydi. Natijada, odamlar turli qarashlardagi axborotlarni kamroq ko'rib, o'ziga o'xshash fikrlarga ega bo'lgan doira ichida qolish ehtimoli oshadi.

Ushbu maqolaning tadqiqot qismida Nordik universiteti talabalari orasida o'tkazilgan nazariy va eksperimental ishlar tasvirlanadi. 1-, 2- va 3-kursdagi jami 100 nafar talaba sinovdan o'tkazildi. Eksperimental usulda ikki yoki uchta guruh tashkil qilinadi: [6]

- **1-guruh:** neytral axborot oqimiga ega bo'ladi (algoritmlar tomonidan filtrlashsiz, turli manbalardan kelgan aralash kontent).
- **2-guruh:** algoritmik tarzda saralangan axborotni oladi (masalan, faqat ma'lum siyosiy yo'nalishdagi yoki biror ijtimoiy mavzu bilan bog'liq kontent).

So'ngra har bir guruhning ishtirokchilari ularga ko'rsatilgan axborotni qabul qilish va baholashga oid bir xil savollarga javob beradi. Bu jarayonda talabalarning kognitiv og'ish darajasi o'lchanadi. Xususan, eksperiment davomida quyidagi kognitiv tarafkashlikka oid o'lchovlar qo'llanadi:

- Confirmation bias: talabalarning o'z e'tiqodlarini tasdiqlovchi ma'lumotlarni tanlash tendentsiyasi.
- Anchoring effect: talabalarning birinchi taqdim etilgan ma'lumot (anchor) ta'sirida qolish darajasi.
- Framing effect: axborot taqdim etilish uslubiga ko'ra qarorlarining o'zgarishi darajasi.

Bu o'lchovlar yordamida guruhlar orasidagi farqlar statistik tahlil qilinadi. Ma'lumotlar tahlili natijasida algoritm yordamida saralangan guruh bilan neytral oqim olgan guruhning javoblari o'rtasida qanday farqlar borligi aniqlanadi.

Ishtirokchilarning ongli va ongsiz tarafkashlik darajasini baholash uchun psixologik metodikalar qo'llanadi. Xususan:

- Cognitive Reflection Test (CRT): bu metodika talabaning intuitiv javob berish o'rniga savolni mulohaza qilib yechish qobiliyatini o'lchaydi aeaweb.org. CRT-da ishtirokchiga oddiy uchta savol beriladi, ular uchida birinchi qarashda «to'g'ri tuyulgan» javob noto'g'ri bo'lib, to'g'ri yechim mantiqiy tahlil talab etadi. Natijada, bu test talabalar orasida intuitiv fikrlash bilan analitik fikrlash darajasini farqlashga yordam beradi.
- Framing Bias Scale (Stanovich, 2018): ma'lumot taqdimoti shakli odamning qarorlariga qanday ta'sir qilishini baholashga mo'ljallangan. Turli kontekstlarda bir xil axborotning prezentatsiya talqiniga qarab, odamlar qarorining o'zgarishini o'lchash uchun ishlatiladi.
- Media Literacy Scale: axborot manbalariga tanqidiy munosabat darajasini o'lchaydi. Ushbu test yordamida talaba algoritmik filtrlarga qanchalik

tanqidiy yondashishini bilish maqsad qilib qo'yiladi. Media savodxonligi baland shaxs algoritm tavsiyalaridan kelib chiqadigan noxolisliklarga kamroq duch keladi deb hisoblanadi.

Ushbu metodikalarning kombinatsiyasi yordamida talabalarning turli tarafkashlik darajalari aniqlanadi. Masalan, CRT talabani dastlabki intuitiv xulosalar bilan solishtirganda haqiqiy fikr yuritishini o'lchaydi, Framing Bias Scale esa ramka effekti natijasida qarorlar qanday o'zgarganini ko'rsatadi [6].

Kontent tahlili (Content Analysis)

Eksperiment doirasida an'anaviy usullarga qo'shimcha tarzda kontent tahlili ham o'tkaziladi. Buning uchun Google yoki YouTube qidiruvlaridan foydalaniladi. Masalan, "vaccines", "religion", "climate change" kabi turli so'zlar bo'yicha qidiruvlar amalga oshiriladi va chiqqan natijalar algoritmik tartibda tahlil qilinadi. Har bir topilgan natijaning:

- Manbasi: (axborot qayerdan olingan – yangilik sayti, blog, ijtimoiy tarmoq va hokazo);
- Reklama holati: natija organikmi yoki reklama orqali efirga uzatilganmi;
- Siyosiy yo'nalishi: agar mavjud bo'lsa, saylovlar, qarashlar, tarafkashlik kuzatiladimi;
- Emotsional (hissiy) tovushi: kontentning ijobiy, salbiy yoki neytral ohangda yozilganligi.

har biri kodlanib chiqiladi. Ushbu ma'lumotlar asosida algoritmning kontentni qanday "ramka"da (taqdim etish shaklida) ajratishi va taqdim etishi baholanadi. Ya'ni, turli ramkalarda berilgan xuddi shu mavzudagi ma'lumotlar qaysi titullar bilan chiqadi va foydalanuvchining qarashlari qanday shakllanishiga olib kelishi o'rganiladi.

Kuzatuv (Observation) va intervyu

Talabalar internetda axborot qidirish jarayonida kuzatiladi. Buning uchun tabiiy muhitda (masalan, laboratoriyada yoki uy sharoitida) foydalanuvchilar biror mavzu bo'yicha ma'lumot qidirayotganda ularning xatti-harakatlari qayd qilinadi. Qidiruvda qidiruv tizimiga kiritilgan so'zlar, bosilgan havolalar, saralangan natijalar bilan qanday ishlagani o'rganiladi.

So'ng, har bir ishtirokchi bilan qisqa intervyu o'tkaziladi. Masalan: "Nega aynan shu manbaga ishonchingiz komil?" yoki "Boshqa manbalarni ko'rib chiqdingizmi?" kabi savollar beriladi. Bu bosqichda ishtirokchilar o'z qarorlarini asoslab, nima uchun ma'lum bir kontentga ishonishni tanlagani, boshqa alternativlarni ko'rib chiqmagani haqida tushuntirish beradi. Bu metod yordamida qatnashchilarning ongli va ongsiz qaror qabul qilish jarayoni, kognitiv tarafkashliklar real va tabiiy holatda qanday namoyon bo'lishi aniqlanadi [5].

Ma'lumotlarni tahlil qilish

Olingan eksperiment va so'rovnomalar ma'lumotlari statistik usullar bilan tahlil qilinadi. Analiz uchun SPSS yoki R dasturlaridan foydalanish rejalashtirilgan. Guruhlar o'rtasidagi farqni aniqlash uchun t-test, ANOVA yoki Chi-kvadrat testlari qo'llanadi. Shuningdek, algoritmik ta'sir bilan har bir o'lchangan kognitiv tarafkashlik darajasi o'rtasidagi bog'liqlikni aniqlash maqsadida korelatsiya va regressiya tahlili ham o'tkaziladi.

Matnli javoblardagi fikrlarni chuqur semantic tahlil qilish uchun NVivo yoki ATLAS.ti kabi dasturlar qo'llanadi. Ular yordamida qatnashchilarning javoblaridagi asosiy mavzular, axborot ramkalari, emotsional tono aniqlanadi va kodlanadi. Bu usul

talabalarning nafaqat rasmiy javoblari, balki intervyudagi so'zlari va kutilmagan tafakkuri asosida xulosa chiqarishga yordam beradi [5].

Kutilyotgan natijalar

O'rganish natijalari shuni ko'rsatadiki, algoritmik tavsiyalar va ijtimoiy tarmoqlarning personalizatsiyasi confirmation bias kabi tarfkashliklarni kuchaytirishi kutilmoqda. Ya'ni, talabalarga algoritmlar tomonidan moslashtirilgan kontent taqdim etilganda ular o'zining asl e'tiqodlarini tasdiqlovchi ma'lumotlarga yanada yopishishi mumkin. Bu esa tanlov erkinligini kamaytiradi – foydalanuvchi faqat o'ziga o'xshagan fikrlarning ichida qoladi va „axborot pufagi“ holati yuzaga keladi.

Shuningdek, «neytral» deb hisoblangan tizimlar ham aslida axborot ierarxiasini belgilash xususiyatiga ega ekani aniqlandi. Masalan, dastlab neytral ko'ringan qidiruv natijalari ham algoritmlar ta'sirida o'z navbatida tartiblanib, ayrim mavzularni yuqori, boshqalarni esa pastroq o'ringa qo'yishi mumkin. Natijada, foydalanuvchi ixtiyorida bo'lgan erkin tanlov maydoni cheklanadi va unga bir tomonlama moslashtirilgan ma'lumot oqimida qolib ketish xavfi paydo bo'ladi.

Xulosa qilib aytganda, vaqtinchalik kuzatuv va laboratoriya tadqiqotlari shuni ko'rsatadiki, algoritmik tafsiv odamlarning qarashlarini toraytirishda muhim rol o'ynaydi. Bu holat kognitiv tarfkashliklar bilan uyg'unlashganida, masalan, tasdiqlovchi tarfkashlik kuchayib, odamlar boshqacha fikrlarni qabul qilishga moyil bo'lmaydi. Natijada, axborot maydoni torayadi va individual qaror qabul qilish erkinligi cheklanadi.

XULOSA

Zamonaviy raqamli muhitda algoritmik tizimlar inson tafakkuri, qaror qabul qilish jarayoni va dunyoqarash shakllanishida markaziy o'rinni egallamoqda. Ushbu maqolada olib borilgan nazariy tahlil va empirik tadqiqotlar shuni ko'rsatdiki, sun'iy intellekt asosidagi algoritmik tafsiv mexanizmlari inson ongida mavjud bo'lgan tabiiy **kognitiv tarfkashliklarni** (confirmation bias, anchoring effect, framing effect) kuchaytiradi. Axborot oqimining shaxsiylashtirilgan tarzda taqdim etilishi foydalanuvchining tanlov erkinligini kamaytirib, uni o'ziga ma'qul bo'lgan fikrlar va qarashlar doirasiga qamab qo'yadi.

Nordik universiteti talabalarida olib borilgan eksperiment natijalari ham bu nazariy xulosalarni tasdiqladi. Neytral axborot oqimini olgan guruhga nisbatan algoritmik tarzda filtrlashgan kontentni olgan guruh ishtirokchilari ko'proq tasdiqlovchi og'ishlarni namoyon etdilar: ular o'z qarashlarini tasdiqlovchi ma'lumotlarni afzal ko'rdi, qarama-qarshi fikrlarga esa shubha bilan qaradi. Shuningdek, "birinchi taqdim etilgan ma'lumotga yopishib olish" holati (anchoring effect) bu guruhda ancha kuchli namoyon bo'ldi. Framing effekti tahlillari esa axborot qanday shaklda — ijobiy yoki salbiy ohangda — berilishiga qarab qarorlar tubdan o'zgarishini ko'rsatdi.

FOYDALANILGAN ADABIYOTLAR RO'YXATI

1. Pariser, E. (2011). *The Filter Bubble: What the Internet Is Hiding from You*. Penguin Press.
URL: https://en.wikipedia.org/wiki/Filter_bubble#:~:text=The%20term%20filter%20bubble%20was,implications%20for%20civic%20%2087

2. Sunstein, C. R. (2009). *Republic.com 2.0*. Princeton University Press. URL:<https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2025.1563592/full>
3. Stanovich, K. E., & West, R. F. (2018). *The Rationality Quotient: Toward a Test of Rational Thinking*. MIT Press.). URL:[britannica.com/science/confirmationbias#:~:text=confirmation%20bias%2C%20people's%20tendency%20to,relevant](https://www.britannica.com/science/confirmationbias#:~:text=confirmation%20bias%2C%20people's%20tendency%20to,relevant)
4. Frederick, S. (2005). *Cognitive Reflection and Decision Making*. *Journal of Economic Perspectives*, 19(4), 25–42.
5. Tversky, A., & Kahneman, D. (1974). *Judgment under Uncertainty: Heuristics and Biases*. *Science*, 185(4157), 1124–1131.
6. Kahneman, D. (2011). *Thinking, Fast and Slow*. Farrar, Straus and Giroux.

