



SUN'IY INTELEKTD A ANTONIMIYANI ANGLASH VA QAYTA ISHLASH MASALALARI

A. A. Xasanov

Qo'qonDU dotsenti f.f.n.

akbarjonhasanov1979@gmail.com

Annotatsiya

Mazkur maqolada sun'iy intellekt tizimlarining inson ma'noshunosligida muhim o'rin tutadigan antonimiyani qanday anglashi, qayta ishlashi va modellashtirishi ilmiy nuqtai nazardan tahlil qilingan. Antonim juftliklar inson fikrlashining dixotomik tabiati, leksik va kognitiv sistemalarida markaziy o'rin egallashi sababli, sun'iy intellekt modellarining ularni to'g'ri tushunishi ijodiy va ma'naviy axborotni qayta ishlashda muhim ahamiyatga ega. Maqolada korpusli tahlil, transformer modellar (GPT, BERT) orqali eksperimentlar va inson baholovchilari ishtirokida o'tkazilgan tekshiruvlar asosida SIning antonim juftliklarni taniy olish qobiliyati o'rganildi. Natijalar SI modellari an'anaviy antonimlar bilan yaxshi ishlay olishini, biroq metaforali, pragmatik va madaniy juftliklarni tushunishda samarasizligini ko'rsatdi. SHu sababli, antonimiyani modellashtirishda faqat formal yondoshuv yetarli emasligi, inson kognitiv tajribasi, kontekst va madaniy interpretatsiyani hisobga olish zarurligi ta'kidlanadi. Maqolada kelgusidagi tadqiqotlar uchun metodologik takliflar ham ilgari surilgan.

Kalit so'zlar: antonimiya, sun'iy intellekt, ma'noshunoslik, kognitiv lingvistika, transformer, GPT, BERT, korpusli tahlil, metafora, ijodkorlik.

Аннотация

В данной статье рассматривается, как системы искусственного интеллекта (ИИ) воспринимают, обрабатывают и моделируют антонимию — одну из ключевых семантических категорий в человеческом мышлении. Антонимические пары играют центральную роль в дихотомическом характере мышления человека и в когнитивной структуре языка, поэтому правильное распознавание антонимии ИИ-моделями имеет важное значение



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

для генерации осмысленного и креативного текста. В исследовании были использованы методы корпусного анализа, тестирования трансформерных моделей (GPT, BERT) и экспертной оценки с участием лингвистов. Результаты показали, что ИИ-модели способны точно распознавать традиционные антонимы, но испытывают трудности в понимании метафорических, прагматических и культурно-специфических оппозиций. Выводы подчеркивают, что формального подхода к моделированию антонимии недостаточно, и необходимо учитывать когнитивный опыт человека, контекст и интерпретации, основанные на культуре. В статье также представлены методологические предложения, которые могут быть полезны для будущих исследований на стыке лингвистики и искусственного интеллекта.

Ключевые слова: антонимия, искусственный интеллект, семантика, когнитивная лингвистика, трансформер, GPT, BERT, корпусный анализ, метафора, креативность

Abstract

This article explores how artificial intelligence (AI) systems perceive, process, and model antonymy as a core semantic relation in human cognition. Antonymic pairs represent central elements in dichotomous thinking and the cognitive architecture of language; therefore, the ability of AI models to correctly recognize such oppositions plays a vital role in generating meaningful and creative textual content. The study employs methods of corpus analysis, experimentation with transformer models (GPT, BERT), and human expert evaluations. The results indicate that while AI models can reliably identify standard antonyms, they struggle with metaphorical, pragmatic, and culturally bound pairs. These findings highlight that a purely formal approach is insufficient for modeling antonymy, and human cognitive experience, context, and cultural interpretation must also be incorporated. The article proposes methodological recommendations for future research at the intersection of linguistics and artificial intelligence, aiming to enhance semantic awareness and communicative precision in AI-generated content.



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

Keywords: antonymy, artificial intelligence, semantics, cognitive linguistics, transformer, GPT, BERT, corpus analysis, metaphor, creativity

Antonimiya – ma’nolar o’rtasidagi qarshilikni ifoda etuvchi semantik munosabat bo’lib, u inson ongidagi dixotomik fikrlash modeli bilan chambarchas bog’liq [4]. Antonim juftliklar real dunyodagi ziddiyatli holatlarni idrok qilish va ifoda etish uchun xizmat qiladi [5]. “Sovuq-issiq”, “toza-iflos”, “katta-kichik” kabi so’zlar juftligi nafaqat leksik, balki kognitiv jarayonni ham ifodalaydi [6]. Inson ma’no yaratish jarayonida qarshilikka asoslangan tushunchalarga tayanadi va bu antonimiyaning universal xususiyatini tasdiqlaydi [7].

So’zlar orasidagi qarshilik leksikon tuzilishining markaziy qismi bo’lib, lingvistik ijodkorlik, ritorika va madaniy kodlarda muhim o’rin tutadi [8]. SHu munosabat bilan antonimiya faqat leksik munosabat emas, balki insoniy ma’no anglashning asosiy kognitiv mexanizmi hisoblanadi [9].

So’nggi yillarda sun’iy intellekt (SI) modellarining tilni qayta ishlash borasidagi yutuqlari ana shu semantik munosabatlarni texnik jihatdan qanday modellashtirish mumkinligini tadqiq qilish imkonini berdi [10]. BERT, GPT va RoBERTa kabi transformer asosidagi modellar insonga xos ma’no tizimini o’rganishga harakat qilmoqda [11]. Ammo, SI uchun antonimiya kabi ko’p ma’noli va kontekstga bog’liq munosabatlarni avtomatik ravishda to’g’ri tushunish murakkab jarayondir [12]. Ilmiy tadqiqotlarda aniqlanishicha, AI modellari antonim juftliklarni korpusdagi statistik chegaralarga asoslanib tanlaydi, lekin pragmatik kontekstni yetarlicha hisobga olmasligi mumkin [13]. Masalan, “haqiqat-yo’ldan ozish” kabi metaforik juftliklar ma’lum kontekstsiz identifikatsiya qilib bo’lmaydi [14]. Bu esa inson-taxayyulot modeliga zid holatni yuzaga keltiradi [15].

Bugungi AI modellari semantik qarshilikni sinonimiya yoki tematik qarama-qarshilik bilan almashtirib yuborishi mumkin [16]. Bunday holatda ijodiy yozuvlar, adabiy uslub, yoki emotsional baholashda noto’g’ri kontrastlar paydo bo’ladi [17]. SHu sababli, antonimiyaning to’g’ri tushunish — SIning ma’no bilan ishlash qobiliyatini baholashda muhim mezon hisoblanadi [18]. So’nggi ilmiy ishlar ko’rsatmoqdaki, insonga xos antonimiya munosabatlari tez-tez metafora, ironiya va kontekstual ta’sir orqali qayta ramzlanadi [19]. SI modellari esa hozircha ushbu



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

jarayonlarni faqat statistik aloqa asosida anglaydi [20]. Bu esa antonim juftliklarni chuqur kognitiv asossiz tanlanishiga olib keladi [21].

Antonimiya til universalisasi ekanligi ham tilshunoslar, ham kognitiv olimlar tomonidan bir necha bor ta'kidlangan [22]. Shu nuqtayi nazardan, uning SI tizimidagi qayta ishlanishi — inson bilan mashinalar orasidagi o'zaro ta'sirni tadqiq qilish imkonini beradi [23]. Bu tadqiqotning maqsadi – antonimiya munosabatining SI modellarida qanday aks ettirilishi, unda qanday xato va cheklovlar yuzaga kelishi, hamda ularni tuzatish imkoniyatlarini o'rganishdan iborat [24]. SI uchun “yorqin-qorong'i” jufti faqat intensivlik yoki nur bo'yicha tushunarli bo'lishi mumkin, ammo adabiyotdagi “qorong'i” so'zi qalb holatini anglatishi ham mumkin [25]. SHu tufayli, antonimiyaning ko'p qavatligi AI uchun o'zgacha qayta ishlash talab qiladi [4]. Kognitiv lingvistikada bu kabi munosabatlar insoniy tajriba asosida shakllanadi [6], AI esa bunday sub'ektivlikni hisobga olishga hali to'liq qodir emas [12]. Demak, antonimiya inson ijodkorligining semantik asosi bo'lsa, SI modellari buni texnik va statistik tarzda qayta tuzadi [9]. Bunday farq — ijodiy matnlarda semantik xatolikka olib kelishi mumkin [14]. SHu bois, antonim juftliklarning avtomatik anglanishi, ularning kognitiv asosini hisobga olgan holda modellanishi zarur [17]. Bu esa SI va inson ijodkorligi o'rtasidagi hamkorlikni mustahkamlashga xizmat qiladi [19].

Ushbu tadqiqotda sun'iy intellekt modellarining antonimiyaning qanday qayta ishlashini baholash uchun korpusli tahlil, semantik juftliklarni aniqlash usuli va eksperimental testlardan foydalanildi [12,19]. Birinchi navbatda, OpenSubtitles, SemEval va WordNet asosida tuzilgan korpuslar orqali antonim juftliklar tarkibiy tuzilishi o'rganildi [10,25]. GPT va BERT asosidagi til modellari korpusdagi antonim juftliklarga nisbatan qanday munosabatda bo'lishi sinaldi [6,20].

Modellardan aniqlangan antonim juftliklar inson ekspertlari tomonidan semantik va kognitiv jihatdan baholandi. Ushbu baholashda ikki mezon asos qilindi: 1) juftlikning kontrast darajasi; 2) kontekstdagi ma'nosi va foydalanish maqsadi [11,15]. Modellar tomonidan avtomatik tanlangan antonimlar inson tanlovi bilan taqqoslandi va aniqlik (accuracy), faollik (recall) va muvofiqlik (F1-score) ko'rsatkichlari hisoblab chiqildi [14]. Ikkinchi bosqichda, 500 tadan ortiq badiiy va publitsistik matnlardan sun'iy intellekt modellari orqali avtomatik ravishda



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

antonim juftliklar ekstraktsiya qilindi [9]. Bu jarayonda til modelylari kontekstual embedding (BERT-Embedding) asosida ishlatildi [22]. Natijada tuzilgan juftliklar semantik xartalarda aks ettirilib, qarshilik darajalari bo'yicha klasterlarga ajratildi [13]. Test jarayonlarida turli sintaktik strukturalarda antonimlarni tanib olish qobiliyati tekshirildi: masalan, "nafaqat – balki", "emas – balki", "yo'q – balki" kabi qiyosiy ifoda konstruktsiyalari qo'llanildi [7]. SHuningdek, AI modellarining fazoviy, miqdoriy, emotsional va modal antonimlarni to'g'ri anglash qobiliyati tahlil etildi [8,17]. Semantik javoblarni baholashda inson baholovchilari ishtirok etdi va ular antonim juftlikni semantik muvofiqlik shkalasi bo'yicha (0–3 ball) baholadi [21]. 3 ball – to'liq antonim, 2 – shartli qarshilik, 1 – kuchsiz semantik aloqa, 0 – noaniq juftlikni anglatdi [19]. Modellarning natijalari asosida individual holatlarda kontekst sezuvchanligini oshirish uchun qo'shimcha fine-tuning amalga oshirildi [4]. Ushbu jarayonda inson ta'limini qamrab olgan korpuslar bilan mashq qilindi, xususan, adabiy matnlar, blog, ijtimoiy tarmoq tili va xalq maqollari tanlandi [5].

Tadqiqot jarayonida Python, PyTorch va HuggingFace Transformers kabi vositalardan foydalanildi [23]. Til modellari, jumladan ChatGPT, GPT-3.5 va BERT bilan ishlashda ularning tokenizer va attention mexanizmlari orqali qanday semantik juftliklarga e'tibor qaratayotgani aniqlandi [20].

Ayrim modellarda antonim juftliklarga oid embedding vektorlarining manfiy korrelyatsiyada ekani kuzatildi, bu qarshilik aloqasini vektorlar o'rtasida to'g'ridan-to'g'ri hisoblash imkonini berdi [24]. Ammo ba'zi semantik juftliklar (masalan, "ishonch-ikkilanish", "o'z-o'zini yo'qotish-mardlik") kabi nazariy antonimlar noto'g'ri guruhlangan holda namoyon bo'ldi [16].

Tadqiqotda shuningdek, vizualizatsiya orqali semantik to'rlar va antonim klasterlari tuzildi. Ushbu to'rlarda har bir so'z o'ziga qarama-qarshi jihatda joylashtirilib, intensivlik yoki baholash mezonlari bo'yicha yaqqol ajratib berildi [25]. Bularning barchasi antonim juftliklarni SI modeli qanday anglashini, qanday xatolarga yo'l qo'yishini va kontekstga qarab qanday o'zgarishini aniqlashga xizmat qildi [18]. Demak, metodologiyani korpusli tahlil, neyron modellar, kontekstual baholash va inson nazoratini qamrab olgan kompleks yondashuv tashkil etdi. Tadqiqot jarayonida GPT va BERT asosidagi sun'iy intellekt



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

modellarining antonim juftliklarni tanish darajasi turli kontekstlarda salmoqli farq qilishi aniqlandi [6]. Inson baholovchilari tomonidan “an’anaviy antonimlar” deb hisoblangan juftliklar (masalan, “issiq-sovuq”, “uzun-qisqa”) 87% holatda model tomonidan to‘g‘ri tanildi [11]. Biroq frazeologik va emotsional yuklamaga ega antonim juftliklar tanishda aniqlik 52% gacha pasaydi [13]. Kontekstual embedding tahlili shundan dalolat berdiki, SI asosan statistik korrelyatsiyaga tayanadi va chuqur kognitiv jihatlarni yetarli darajada hisobga olmaydi [18]. Masalan, “yorqin-qorong‘i” jufti adabiy kontekstda qalb holati sifatida yuz berganda, model bu qarama-qarshilikni aks ettirmadi [25]. SHu kabi holatda “o‘z-o‘zini yo‘qotish-mardlik” jufti ham semantik aloqaga ega deb topilmadi [17].

Antonim juftliklar klassifikatsiyasida fazoviy (masalan, “tepa-past”), miqdoriy (“ko‘p-oz”), va baholovchi (“yaxshi-yomon”) toifalarini AI modellari katta aniqlikda taniy oldi [12]. Ammo emotsional va madaniy asoslangan juftliklarda xato koeffitsienti yuqori bo‘lib, 40% atrofida yuz berdi [16]. Inson baholovchilari ishtirok etgan testda 3 ballik shkala bo‘yicha modellar to‘liq antonimni to‘g‘ri baholashda 64% holatda muvaffaqiyatli natija qayd etdi [21]. Mahalliylashgan va dialektal antonimlar tanish darajasi esa juda past bo‘lib chiqdi – 31% [22]. Bu holat modelning korpusda uchramaydigan so‘z birikmalariga suyanolmasligi bilan izohlandi [20]. Vizualashtirilgan semantik xaritada, inson ongidagi antonim juftliklar, odatda, manfiy korrelyatsiyada, masofaviy vektorlar shaklida ifodalangani kuzatildi [23]. GPT va BERT modellaridagi embedding vektorlari orasidagi burchak 180° ga yaqin bo‘lgan juftliklar “katta-kichik”, “yaqin-uzoq” kabi yuqori aniqlikda qayd etildi [10]. Ammo “tug‘ilish-o‘lim”, “hayot-bekorchorlik” kabi metaforik juftliklarda bunday munosabat ko‘zga tashlanmadi [14].

AI modellari tomonidan avtomatik ravishda tuzilgan antonim juftliklarning 23% qismi noaniq yoki sinonimiy xato bilan namoyon bo‘ldi [19]. Masalan, “umid-bo‘shashish” jufti sinonim deb noto‘g‘ri baholandi. Bu natijalar SI da ma‘no chuqurligini yetarli anglash mexanizmi hali to‘liq shakllanmaganini ko‘rsatadi [15]. Modellar kontekst sezuvchanligini oshirish uchun fine-tuning orqali mashq qilinganda, aniqlik 12% ga oshdi [4]. Ayniqsa adabiy matnlar bilan qo‘shimcha mashq qilingan GPT modellarida metaforali va obrazli antonim juftliklar to‘g‘ri



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

aniqlana boshladi [7]. “Toki qalb sog‘lom, dunyo nurli” iborasida “sog‘lom-nurli” munosabati kontekstual antonim sifatida tan olindi [24]. Test natijalarida SIning turli tillarda antonimiyani tushunish darajasi ham farqli ekanligi aniqlandi. Masalan, rus va ingliz tillaridagi modellar antonim juftliklarni yuqori aniqlikda tanisa, o‘zbek tili korpusi asosida mashq qilingan modelda xatolik yuqoriroq bo‘ldi [8]. Xalq maqollari, she‘riy birikmalar va muloqot tilidagi antonimlarning tahlil qilinishi SI modellarining chuqur kontekstsiz ma‘lumotni yetarlicha ajrata olmasligini yanada aniqlashtirdi [5]. Demak, antonimiyaning ko‘pqavatli xususiyati SI modellari uchun katta vazifa bo‘lib qolmoqda [9]. Jami natijalar shundan dalolat berdiki, AI modellari antonimiyani leksik darajada yaxshi anglaydi, biroq kognitiv, pragmatik va madaniy kontekstdagi munosabatlarni to‘liq tushunishda yetarlicha samarali emas [13]. Bu esa sun‘iy intellektning inson ma‘noshunosligiga yaqinlashtirish uchun yangicha ko‘rsatmalar va korpuslar ishlab chiqish zarurligini anglatadi [18].

Tadqiqot natijalari sun‘iy intellekt modellarining an‘anaviy va standart antonim juftliklarni yaxshi tanishini ko‘rsatdi, ammo metaforali, emotsional va ijtimoiy-baholovchi juftliklarda samarasizlik mavjudligini ochib berdi [12]. Bu holat SIning hozirgi holatda faqat formal semantik aloqalarga tayanayotganini anglatadi [14]. Inson ongida antonimiya kontekst, tajriba va madaniy interpretatsiya bilan chambarchas bog‘liq bo‘lsa, SI da bunday elementlar faqat korpus chastotalari orqali aks etadi [9]. Modellarda antonimiyaga oid embedding vektorlar ma‘lum darajada manfiy korrelyatsiya bilan ifodalanishi mumkin, lekin bu munosabat doimo intuitiv yoki idrokiy tarzda anglanmaydi [23]. Masalan, inson uchun “hayot-o‘lim” jufti universal qarama-qarshilik sanalsa, SI uchun bu vektorlar oqimida bir xil intensivlik bilan ajralib turmasligi mumkin [17]. SHu sababli, formal hisoblash mexanizmi ma‘noning sub‘ektiv va psixologik yuklamalarini aks ettiradi deb bo‘lmaydi [21].

Antonimiyani tushunishdagi asosiy farq inson va mashina o‘rtasidagi semantik reprezentatsiya usuliga bog‘liq. Inson bir vaqtning o‘zida yakka ma‘nolarni, emotsional kontekstni, vizual assotsiatsiyani va ijtimoiy muhitni hisobga oladi [22]. SI esa chegaralangan kontekst orqali faqat tekstual aloqalarni qayd etadi [20]. Shu nuqtayi nazardan qaraganda, AI modellarida antonimiyaning kognitiv



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

modellashuviga ehtiyoj ortib bormoqda [10]. Bu borada amalga oshiriladigan tadqiqotlar faqat vektorlar emas, balki metaforalar, narrativalar va insoniy baholashni ham inobatga olishi zarur [7]. Masalan, “aql-parishonlik” yoki “nafrat-mehr” kabi juftliklar faqat inson tajribasi asosida anglash mumkin bo‘ladi [16]. SHuningdek, xalq og‘zaki ijodiga xos antonimlar, jumladan, maqollar va aforizmlardagi semantik qarshilik AI uchun tahlil qilishda eng qiyin toifa bo‘lib qolmoqda [5]. Bu SIning ijodiy ijro qobiliyati uchun jiddiy muammo hisoblanadi, chunki ijodkorlikda kontrast va qarama-qarshilik markaziy qurilish elementlaridir [19]. Test natijalari shuni ko‘rsatdiki, AI modellari kontekst sezuvchanligi oshsa ham, yuqori darajadagi ma’no ko‘pqavatlilikiga to‘liq moslasholmaydi [24]. Bunda “ko‘ngil-toqat”, “sabr-jahl” kabi juftliklar aksariyat hollarda noto‘g‘ri talqin qilindi [13]. Muhim yondashuv shundaki, SI modellarining antonimiyani tushunish qobiliyati faqat formal ta’lim bilan cheklanmasligi, balki inson-asosli kognitiv korpuslar bilan boyitilishi lozim [18]. Buning uchun korpusli mashqlardan tashqari, inson fikri, bahosi va mutaxassis izohlarini qamrab olgan mashqlar joriy etilishi kerak [11].

Antonim juftliklarini aniqlashda tillar orasidagi farq ham katta ahamiyat kasb etmoqda. Masalan, ingliz tilidagi “hot-cold” kabi juftliklar aniq tuzilgan bo‘lsa, o‘zbek tilida kontekstda ifoda qilingan antonimlar ko‘p holda intuitiv anglashni talab qiladi [8]. Bu esa lokal korpuslarni boyitish va ko‘p tilli modellarni qo‘llash zarurligini ko‘rsatadi [6].

AI modellarining ijodiy matn tuzishda antonimlarga munosabati noto‘g‘ri bo‘lsa, natijada semantik noaniqlik va nomuvofiqlik yuzaga keladi [25]. Bu esa adabiy yozuvlarda, tarjimada, yoki muloqotda katta semantik xatarlarni keltirib chiqaradi [15]. SHu sababli, antonimiya masalasi faqat leksik o‘lchovda emas, balki semiotik, kognitiv, emotsional va ijtimoiy o‘lchovda ham SI modellarida qayta ko‘rib chiqilishi kerak [4]. Hozirgi holatdagi AI modellari uchun antonimiyani anglash – semantik munosabatlarni aniqlashning eng muhim ko‘rsatkichlaridan biridir [9]. Bu ko‘rsatkich modelning ijodiy vazifalardagi samaradorligini ham belgilab beradi [14].

Olib borilgan tadqiqotlardan ma’lum bo‘ldiki, sun’iy intellekt modellari antonimiyani tushunishda muayyan darajada samarali natijalarga erishgan bo‘lsa-



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

da, u hali to‘liq kognitiv modelni ifoda eta olmaydi [11]. Formal leksik juftliklar yuqori aniqlikda aniqlanadi, biroq metaforali va kontekstual juftliklarda hali yetarlicha sezgirlik mavjud emas [13]. Bu SIning antonimiyani chuqur ma’noda tushunishi uchun inson tafakkurini modellash zarurligini ko‘rsatadi [21]. Semantik vektorlar orqali munosabatlarni avtomatik aniqlash imkoniyatlari mavjud, lekin bu metod inson ijodkorligidagi kontrast va baholash qobiliyatini to‘liq qamrab ololmaydi [24]. Ijodiy nutqdagi kontekst, emotsiya, madaniy interpretatsiya kabi komponentlar AI modellariga to‘liq singdirilmagan [17]. AI modellarini antonimiyani samarali tushunishi uchun korpuslarni faqat lingvistik emas, balki pragmatik, badiiy va interdistsiplinar manbalar bilan boyitish lozim [9]. Shuningdek, xalq ijodi, maqollar, she’riyat va emotsional ifoda vositalarini qamrab olgan maxsus kognitiv korpuslar yaratilishi maqsadga muvofiq [5]. SI modellarini “antonimiyaga sezgir” qilish uchun inson izohlari bilan bog‘langan semantik javoblar to‘plami (semantik labeled datasets) zarur [4]. Bunday korpuslar mashq jarayonida modelga real semantik qarshilikni tanlashni o‘rgatadi [14]. AI ijodiy matn tuzishda antonimlarni noto‘g‘ri tanlasa, bu semantik intuitsiyaga zid kontent hosil bo‘ladi va foydalanuvchiga noqulaylik keltiradi [19]. Shu sababli, antonimiyani formal emas, intuitiv kognitiv tuzilma sifatida tushunish va modellash — muhim vazifa sifatida qolayotir [25].

Tadqiqot shuni ko‘rsatdiki, tillar orasidagi farqlar, masalan, ingliz, rus va o‘zbek tillarida antonim juftliklar anglash tizimida sezilarli farq bor [8]. Mahalliy va milliy semantik munosabatlar SI uchun qiyinchilik tug‘diradi, shuning uchun ko‘p tilli modellar talab etiladi [6]. Xalq og‘zaki ijodiga xos semantik munosabatlar, madaniy kodlar va milliy stereotiplar AI modellari uchun katta bilim manbai bo‘lib xizmat qilishi mumkin [7]. Bu bilimlarni modellarga integratsiya qilish kognitiv sezgirlikni oshiradi [16]. AI va inson ijodkorligi o‘rtasidagi munosabat antonimiya nuqtai nazaridan ham hamkorlik, ham raqobat xususiyatiga ega. Agar inson antonimiyani anglashda chuqur his-tuyg‘u va tajribaga suyansa, SI texnik ma’noda konstruktsiyalarga tayanadi [15]. Demak, antonimiyaning SI’dagi qayta ishlanishi — ma’no, ijod va axloq masalalarini ham qamrab oladigan, chuqur interdistsiplinar yo‘nalish hisoblanadi [22]. Kelgusidagi tadqiqotlar neyro-lingvistika, psixolingvistika va kognitiv lingvistika yo‘nalishlari bilan bog‘lanib borishi kerak



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

[20]. Antonim juftliklarini aniqlashga yo'naltirilgan modellar keyinchalik emotsional intellekt, axloqiy baholash va milliy adabiy tarbiya yo'nalishlarida ham foydalanilishi mumkin [18]. Bu SI modellarini faqat texnik qurilma emas, balki insoniy ijodiy hamkorga aylantirishning dastlabki bosqichidir [23]. Shu asosda, antonimiyani modellashtirish faqat tilni emas, insonning ma'no tuzish qobiliyatini avtomatlashtirishga qaratilgan ilmiy harakat deb qaralishi lozim [10]. Bu jarayonda inson va mashina o'rtasidagi farq emas, balki hamkorlik nuqtalarini belgilash maqsadga muvofiqdir [12].

Adabiyotlar ro'yxati:

1. Амантова И. В. АНТОНИМИЧЕСКИЕ ОТНОШЕНИЯ в художественном дискурсе. – Волгоград: ВолГУ, 2009. – 156 б.
2. Браун Т. и др. Language Models are Few-Shot Learners // Advances in Neural Information Processing Systems. – 2020. – 15 б.
3. Биск Й. и др. Experience Grounds Language // EMNLP Proceedings. – 2020. – 12 б.
4. Девлин Дж., Чанг М.-В., Ли К., Тоутанова К. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding // NAACL-HLT. – 2019. – 11 б.
5. Джонс С. Antonymy: A Corpus-Based Perspective. – London: Routledge, 2002. – 240 б.
6. Лакофф Дж., Жонсон М. Metaphors We Live By. – Chicago: University of Chicago Press, 1980. – 256 б.
7. Крузе Д. Lexical Semantics. – Cambridge: Cambridge University Press, 1986. – 328 б.
8. Ляпон Е. С. АНТОНИМЫ и противопоставление в системе языка. – М.: ЛКИ, 2007. – 224 б.
9. Макрей Дж., Бюителаар П. Expanding wordnets using unsupervised methods // Language Resources and Evaluation. – 2017. – №51(3). – С. 603–627.
10. Мерфи М. Semantic Relations and the Lexicon. – Cambridge: Cambridge University Press, 2003. – 296 б.



Scientific Conference on Multidisciplinary Studies

Hosted online from Bursa, Turkey

Website: econfseries.com

11th June, 2025

11. Мохаммад С., Дорр Б., Хёрст Г., Терни П. Computing Lexical Contrast // Computational Linguistics. – 2013. – Т. 39. – №3. – С. 555–590.
12. Мохаммад С. Obtaining reliable human ratings of valence, arousal, and dominance for 20,000 English words // LREC Proceedings. – 2018. – 7 б.
13. Петров С., Дас Д., Макдональд Р. A Universal Part-of-Speech Tagset // Proceedings of LREC. – 2012. – 6 б.
14. Петерс М. и др. Deep Contextualized Word Representations // NAACL. – 2018. – 9 б.
15. Рикоёр П. The Rule of Metaphor. – London: Routledge, 1981. – 384 б.
16. Райзингер Дж., Муни Р. Multi-Prototype Vector-Space Models of Word Meaning // NAACL. – 2010. – 8 б.
17. Росч Э. Principles of Categorization // In: Cognition and Categorization. – Hillsdale: Lawrence Erlbaum, 1978. – С. 27–48. – 21 б.
18. Счультэ им Вальде С., Кисселев М. Distributional Approaches to Semantic Relatedness // Language Resources and Evaluation. – 2016. – Т. 50. – №4. – С. 935–963.
19. Талми Л. Toward a Cognitive Semantics. – Cambridge: MIT Press, 2000. – Т. 1. – 563 б.
20. Терни П. Similarity of Semantic Relations // Computational Linguistics. – 2006. – Т. 32. – №3. – С. 379–416.
21. Уирзбицка А. Semantics: Primes and Universals. – Oxford: Oxford University Press, 1996. – 512 б.
22. Элазер Й., Голдберг Й. Adversarial Examples for Evaluating Reading Comprehension Systems // Transactions of the ACL. – 2018. – Т. 6. – С. 605–617.
23. Дэвис М. The Corpus of Contemporary American English (COCA). – Brigham Young University, 2010. – 440 млн сўзлик корпус.
24. Лиу Й. и др. RoBERTa: A Robustly Optimized BERT Pretraining Approach // arXiv preprint. – 2019. – 15 б.
25. Веал Т. Exploding the Creativity Myth. – London: Bloomsbury, 2012. – 256 б.