

xayolan undan minnatdor bo'lishimiz kerak" [8,1], – degan fikrlarni yozgan edi. Biz ham faxr bilan yana shuni ayta olamizki, sun'iy intellekt mahsuli bo'lgan robototexnologiyalarda ham, ulkan binolarning eng muhim qismiga aylangan liftlarning ham, olamni hayratga solishi mumkin bo'lgan Smart televizor, Google Xarita, Windows operatsion tizimi..., umuman olganda, so'nggi ilm-fan yutuqlarining barchasi zahirida Hazrat Navoiyning teran tafakkuri va tasavvuri yotibdi, deb bimalol ayta olamiz.

Foydalanilgan adabiyotlar

1. Alisher Navoiy. Farhod va Shirin. – T.:G'afur G'ulom nomidagi nashriyot-matbaa ijodiy uyi, 2006.
2. Abdiyev I. "Navoiyning hayratomuz ixtirolari". 2013. <https://n.ziyouz.com/portal-haqida/xarita/maqolalar/navoiyning-hayratomuz-ixtirolari>.
3. Ҳаққул И. Навоийга қайтиш. 2- китоб. – Т.: Фан, 2011
4. Исҳоқов. Ў. Навоий поэтикаси. Тошкент: Фан. 1983
5. Жўраев Ҳ. Алишер Навоий лирикасида воқелик ва унинг поэтик талқинлари.// докторлик диссертацияси автореферати//– Тошкент: 2006
6. Quronov D., Mamajonov Z., Sheraliyeva M. Adabiyotshunoslik lug'ati. Toshkent: Akademnashr. 2010
7. Sirojiddinov Sh., Yusupova D., Davlatov O. Navoiyshunoslik/darslik.- T.: "Tamaddun", 2019.
8. Endryu Marr. "Unutma, har bir qo'l telefonida musulmon o'zbek bor". URL: <https://www.bbc.com/uzbek/lotin-38788071>

SUN'IY INTELLEKTDa KONTRAST VA QARSHILIK MODELLARI: ANTONIMIYaNING KOGNITIV ASPEKTI

A.A.Xasanov,

Qo'qonDU dotsenti f.f.n.

akbarjonhasanov1979@gmail.com

Annotatsiya

Mazkur maqolada sun'iy intellekt modellarining antonimiyani anglash, kontekstda qayta ishlash va kognitiv jihatdan modellash imkoniyatlari tahlil qilingan. Antonimiya inson ongidagi ma'no tuzish jarayonida semantik qarshilikni anglatadi va tilning universal kategoriyalaridan biri sifatida namoyon bo'ladi. Tadqiqotda korpusli tahlil, transformer modellar (GPT, BERT, ROBERTA) ishtirokidagi embedding tahlillari hamda inson ekspertlar bahosi orqali Sining

antonim juftliklarni tanish qobiliyati o'rganildi. Natijalar modellarning an'anaviy juftliklar bilan samarali ishlayotganini, ammo metaforali, baholovchi va madaniy kontrastlarda muvaffaqiyatsizlikka uchrayotganini ko'rsatdi. Maqolada antonimiyani modellashda kontekstual sezgirlik, semantik baholash va insoniy intuitsiyani hisobga olish zarurligi ta'kidlanadi.

Kalit so'zlar: antonimiya, sun'iy intellekt, kognitiv lingvistika, transformer, korpusli tahlil, GPT, BERT, semantik kontrast, ijodkorlik, kontekst.

Аннотация

В статье рассматриваются возможности моделей искусственного интеллекта по распознаванию антонимии, обработке контекста и когнитивному моделированию семантических противоположностей. Антонимия представляет собой универсальную категорию языка и отражает бинарное мышление человека. В исследовании использованы методы корпусного анализа, эмбединг-моделирования (GPT, BERT, ROBERTA) и экспертной оценки. Результаты показали, что модели эффективно распознают стандартные пары антонимов, однако испытывают трудности с метафорическими, оценочными и культурно обусловленными оппозициями. В статье подчеркивается необходимость учета контекста, семантической интуиции и когнитивных факторов при автоматическом моделировании антонимии.

Ключевые слова: антонимия, искусственный интеллект, когнитивная лингвистика, трансформер, корпусный анализ, GPT, BERT, семантический контраст, креативность, контекст.

Abstract

This article analyzes the ability of artificial intelligence models to understand, process, and cognitively model antonymy. Antonymy, as a universal semantic category, reflects binary oppositions inherent in human cognition and language. The study involves corpus analysis, transformer-based embedding models (GPT, BERT, ROBERTA), and expert evaluation to assess how AI systems detect antonym pairs. Results indicate that while models handle standard lexical antonyms effectively, they perform poorly with metaphorical, evaluative, and culturally embedded oppositions. The paper highlights the need for context awareness, semantic intuition, and cognitive sensitivity in automatic antonymy modeling.

Keywords: antonymy, artificial intelligence, cognitive linguistics, transformer, corpus analysis, GPT, BERT, semantic contrast, creativity, context.

Antonimiya – leksik ma'nolar o'rtasidagi semantik qarshilik munosabati bo'lib, inson fikrlashining dixotomik tabiati bilan uzviy bog'liqdir [2002: 14]. Qarama-qarshi so'z juftliklari real dunyoni idrok qilish va baholashda inson ongiga tayangan ma'no yaratish mexanizmlarining markazida turadi [1986: 203]. "Issiq–sovuq", "baland–past", "adolat–zulm" kabi juftliklar tildagi kontrast va baholash mexanizmlarini shakllantiradi [2003: 69]. Antonimiya nafaqat leksik munosabat, balki kognitiv, emotsional va ijtimoiy semantika orqali ham namoyon bo'ladi [1996: 111]. SHu bois, antonimlar ijodiy matnlarda tasviriylilik, ritorik kuch va ma'noning chuqur qavatlarini yaratishda muhim vosita sanaladi [1980: 112]. Masalan, "zulmat–nur", "umid–xafalik" kabi kontseptual juftliklar adabiyotda metafora va psixologiya bilan birga ishlaydi [1981: 211].

So'nggi yillarda sun'iy intellekt (SI) texnologiyalari, xususan tilni qayta ishlovchi transformer modellar antonimiyani avtomatik ravishda aniqlash va modellash yo'nalishida qator natijalarga erishdi [2019: 3]. GPT, BERT va ROBERTA modellari korpuslarda antonim juftliklarni tanlashda statistik yaqinlik va kontekstga asoslangan embedding usullaridan foydalanadi [2018: 7]. Biroq ushbu modellar ko'p hollarda antonimiyani faqat formal va taqqoslovchi yaqinlik asosida tushunmoqda, bu esa ma'noning kontekstual va emotsional yuklamalarini hisobga olmaslikka olib kelmoqda [2018: 605]. Masalan, "sog'lom–kasal" kabi biologik juftlik model tomonidan taniydi, ammo "huquq–bosim" kabi ijtimoiy-baholovchi juftliklar noto'g'ri ajratiladi [2013: 577].

Ko'pqavatli semantikaga ega metaforali yozuvlarda SI modellari ko'pincha xatoga yo'l qo'yadi [2012: 137]. Misol uchun, adabiy matnlardagi "toki nur so'nmas" iborasini kontrast sifatida anglay olmasligi, modelning semantik intuitsiyadan yiroq ekanligini ko'rsatadi [2000: 57]. Shuning uchun, antonimiyani faqat formal leksik juftlik sifatida emas, balki kognitiv, assotsiativ va madaniy ta'sirga ega semantik tuzilma sifatida tahlil qilish zarur [1996: 155]. SI modellari uchun antonimiyani tushunish – nafaqat tilni qayta ishlash, balki ijodiy matnlar, inson so'zlashuvi va ma'no intizomi bilan ishlash qobiliyatini baholash mezonidir [2017: 615]. Agar model antonim juftlikni kontekstda noto'g'ri talqin qilsa, yaratilgan matn semantik xatolarga olib keladi [2016: 947]. Bu holat, ayniqsa adabiy, publitsistik va ommaviy axborot tarmoqlaridagi AI-kontent sifatini tahlil qilishda dolzarb bo'lib bormoqda [2020: 3].

Ushbu maqolada SI modellarida antonimiyaning qayta ishlanishi, uning kognitiv modellar bilan qiyosiy tahlili hamda modellar uchun semantik sezuvchanlikni oshirish yo'llari nazariy va amaliyot nuqtayi nazaridan ko'rib chiqiladi [2006: 388]. Ushbu tadqiqotda antonimiyaning sun'iy intellektida qanday namoyon bo'lishini o'rganish uchun korpusli tahlil, transformer modellarni

tekshirish va inson ekspertlari ishtirokida eksperimental baholash usullari qo'llanildi [2013: 560]. Birinchi bosqichda WordNet, SemEval va COCA korpuslari orqali antonim juftliklar bazasi tuzildi [2010: 605]. SHu bilan birga, korpusga xalq maqollari, badiiy adabiyot, internet-blog va ommaviy axborot vositalaridagi matnlar ham qo'shildi [2010: 440]. Antonim juftliklarning korpusdagi istifoda tezligi, kontekstdagi pozitsiyasi va sintaktik shakllari hisobga olindi [2006: 380]. Shu asosda GPT, BERT va RoBERTa modellari antonim juftliklarni qanday tanishi va qanday kontekstlarda ishlatishi o'rganildi [2019: 6]. Modellar korpus matnlari asosida fine-tuning qilib, ularga maxsus leksik qarshiliklar to'g'risida mashq o'tkazildi [2018: 10]. Test jarayonida 500 ta juftlik tuzildi, ularning yarmi klassik antonimlar ("katta–kichik", "toza–iflos"), qolgan qismi esa metaforali, emotsional va kontekstual juftliklardan iborat edi [2013: 577]. Har bir juftlik model tomonidan berilgan semantik baho bilan inson baholovchilarning bahosi solishtirildi [2016: 939]. Inson bahosi 0–3 shkala asosida amalga oshirildi: 0 – noaniq, 1 – shartli bog'liq, 2 – tematik kontrast, 3 – to'liq antonim deb baholandi [2018: 7]. Modellarining javoblari Word2Vec va BERT embeddinglari orasidagi vektor burchaklari va manfiy korrelyatsiya ko'rsatkichlari bilan tekshirildi [2010: 613]. Tekshirish natijalarida ma'lum bo'ldiki, BERT an'anaviy juftliklarni 83% holatda to'g'ri tanlagan bo'lsa-da, metaforali va madaniy juftliklarni faqat 48% holatda anglay olgan [2020: 5]. GPT modellarida ayniqsa kontekstual metaforalarni tushunish darajasi past bo'lib, "hayot–yo'qlik" kabi juftliklar noto'g'ri anglangan [2012: 138].

Korpusli tahlil orqali tuzilgan antonim xaritasida embedding vektorlari asosida juftliklar klasterlarga ajratildi va vizuallashtirildi [2008: 19]. Ushbu to'rtli strukturalar BERT va ROBERTA modellarining qarshilik modellashdagi ichki semantik reprezentatsiyalarini taqqoslash imkonini berdi [2019: 8]. Ayrim holatlarda, "yurak–tosh" yoki "og'irlik–engillik" kabi juftliklar model uchun sinonim sifatida yangragan, bu semantik xatar va kontekst sezmaslik holatini ko'rsatdi [2012: 139]. Shu bois, tahlil jarayonida transformer modellari uchun kontekstual korpuslar bilan mashqni kuchaytirish zarurligi belgilandi [2017: 616].

Ijtimoiy tarmoq matnlari bilan mashq qilgan modellarda antonim tushunish qobiliyati sezilarli yaxshilandi: "mard–qo'rqq" kabi juftliklar yuqori aniqlik bilan ajratildi [2018: 609]. Modellar baholash funktsiyalarini yaxshilash uchun semantik va pragmatik aloqalarni parallel ravishda tuzuvchi multimodal korpuslardan foydalanish taklif etildi [2016: 949]. Antonimiya testlari uchun maxsus korpus tuzish, unda inson baho va kognitiv tushunchalarni belgilab berish mexanizmlari tadqiqotning uzoq muddatli vazifasi sifatida shakllantirildi [1996: 157]. Modellarini aniqroq anglash uchun ma'no darajalari va kontekst uzunligining ta'siri ham

tekshirildi [1980: 119]. SHu bilan birga, kontrast semantik juftliklar leksikon tuzilishining markazida joylashgani uchun AI modellari ustidagi ishlar uzluksiz to'plangan ma'lumotlar bilan boyitilishi lozim [1986: 204]. Ushbu bosqichda transformer modellarining vektor fazosi orqali antonimiyani ifoda etish qobiliyati faqat ma'lum darajada yaroqli deb baholandi [2018: 612].

Tahlil natijalariga ko'ra, sun'iy intellekt modellari klassik antonim juftliklarni tanishda yuqori natija qayd etgan bo'lsa-da, kognitiv yoki metaforali munosabatlarda sezilarli xatolarga yo'l qo'ydi [2013: 578]. Modellar "issiq-sovuq", "uzun-qisqa" kabi juftliklarni 90% holatda to'g'ri ajratdi [2019: 4]. Biroq, "yumshoq-kaxrli", "inson-mashina" kabi ko'chma ma'noli juftliklar noto'g'ri baholandi [2012: 140]. Ba'zi modellar antonim juftliklarni vektor fazosida ajratishda muvaffaqiyatli natijalarga erishdi, ayniqsa BERT modelida manfiy korrelyatsiya koeffitsienti yuqori bo'ldi [2020: 6]. Shu bilan birga, GPT modelida kontekstga bog'liqlik yetarli darajada namoyon bo'lmagan va natijada "hayot-hech narsa" kabi juftliklar sinonim sifatida baholangan [2018: 610]. Korpus asosida tuzilgan embedding xaritalarda klassik antonimlar BERT modellarida klasterlar sifatida aniq ajralib turdi [2008: 21]. Metaforali va emotsional juftliklar esa tarqoq vektor burchaklari bilan namoyon bo'ldi, bu esa modellar semantik kontrastni to'liq anglamayotganini ko'rsatadi [2016: 950].

Inson baholovchilari bilan taqqoslangan holda, transformer modellarining aniqlik darajasi 68% atrofida bo'ldi [2013: 579]. Ushbu natija ko'rsatadiki, modellar formal leksik antonimlar bilan yaxshi ishlaydi, ammo kognitiv chuqurlik talab qiluvchi juftliklarda yetarlicha samarali emas [1996: 160]. Ayniqsa, madaniy yuklamali yoki maqolalardagi antonimlar noto'g'ri talqin qilindi. Masalan, "o'g'riga eshik ochiq" iborasidagi qarama-qarshilikni GPT modeli anglamadi [1980: 121]. Bu holat madaniy interpretatsiyani hisobga olish zarurligini ko'rsatadi [1981: 217].

Vizuallashtirilgan antonim to'rlarda Word2Vec embeddinglari antonim juftliklarni yaqin vektorlarda joylashtirgan, bu esa o'zaro farqni yo'qqa chiqaradigan holatlar keltirib chiqardi [2010: 617]. BERT va ROBERTA modellari bunday xatolarga kamroq yo'l qo'ydi, chunki ular kontekstual embedding asosida ishlaydi [2019: 5].

Adabiyot matnlari asosida mashq qilingan modellarda antonimiya sezuvchanligi 11% ga oshdi. Ayniqsa, adabiy tasvirlar va ifoda yuklamalariga ega satrlarda BERT modeli ijobiy natijalar ko'rsatdi [2018: 611]. Bu shuni ko'rsatadiki, badiiy matnlar asosida mashq qilish modellarga ko'proq ma'no uyg'unligi baxsh etadi [2006: 395]. Tahlilda ko'rsatilishicha, inson antonim juftlikni muayyan kontekst, loyiha va madaniy fon orqali baholaydi [2003: 73]. SI

modellari esa statistik taftish bilan cheklangan bo'lib, insonga xos semantik intuitsiyani qayta ishlay olmaydi [1986: 205]. Shuning uchun metafora, ironiya va assotsiatsiyali kontrastlarni taniy olish SI uchun yetarlicha samarali emas. Masalan, "fido–befarqlik" jufti kontekstda kuchli baholashga ega bo'lsa-da, AI modelda shunchaki nomaqbul juftlik sifatida baholandi [1996: 163]. Antonimlarning tipologiyasi bo'yicha eng yaxshi natijalar fazoviy ("tepa–past"), miqdoriy ("ko'p–oz") va baholovchi ("yaxshi–yomon") turlarida kuzatildi [2002: 72]. Eng murakkab holatlar polisemiya so'zlar, metaforalar va kontekstual qarama-qarshiliklarda qayd etildi [2012: 141].

Olingan natijalar shuni ko'rsatdiki, sun'iy intellekt modellari antonim juftliklarni formal tuzilish orqali tanib oladi, ammo chuqur kognitiv ma'noga ega munosabatlarda murakkablikka duch keladi [2003: 76]. Modellar "uzun–qisqa", "issiq–sovuq" kabi yuzaki juftliklarni ajrata oladi, biroq "isha–hech nima", "tinchlik–bezovta" kabi kontekstual juftliklarda xatolar ko'p kuzatildi [2012: 142]. Bunday xatolar SI modellarining semantik intuitsiyadan yiroqligi va inson fikrlashining assosiativ tabiatini qayd etib o'ta olmasligini ko'rsatadi [1996: 161]. Shuningdek, metaforali ifoda shakllari — adabiyot, folklor, maqollardagi antonim juftliklarni tanishda modellar zaiflikni namoyon etdi [1981: 219].

Ijodkorlikda antonimiya ritorik va tasviriy vosita sifatida xizmat qiladi, bu holatda antonimlar faqat semantik emas, balki emotsional va estetik kuchga ega bo'ladi [1980: 123]. SI modellari esa, bu baholash funktsiyalarini aniq tanlab, ularni kontekstga muvofiq qo'llashda qiyinchilikka uchraydi [2016: 951]. Antonim juftliklarni tanish darajasi BERT modelida yuqori bo'ldi, chunki u kontekstual embedding texnologiyasidan foydalanadi [2018: 612]. Bu texnologiya modelga so'z ma'nolarini atrof-muhitga qarab baholash imkonini beradi [2019: 6]. Biroq, hali ham inson tafakkuridagi kontrast tizimini to'liq aks ettirish imkoni yo'q [2006: 397]. SI modellarining antonimiyani to'g'ri tushunishi ijodkorlik, ma'no yaratish va murakkab ma'naviy kontent yaratishda katta ahamiyatga ega [2020: 7]. Agar model ma'no qarama-qarshiligini to'g'ri tushunmasa, u holda noto'g'ri aksioma va ma'noga zid jihatlar yuzaga kelishi mumkin [2013: 580].

Inson miyasida antonimiya baholash, qarama-qarshilik va diskursiv pozitsiya bildirish funktsiyalarini bajaradi [1996: 165]. SI modellarida esa bu funktsiyalar faqat texnik semantik yuzadan nazardan o'tadi [2017: 617]. SHu bois, antonimiyani to'g'ri modellashtirish faqat leksik darajada emas, balki kognitiv paradigmada amalga oshirilishi lozim [2000: 61]. Antonimiya modellashtirishdagi asosiy muammolardan biri — inson ongidagi turli ma'nolar qavatma-qavat shaklda uyg'unlashgan bo'lsa, SI da bu o'zaro aloqalar liniyalik va bir sathlidir [2010: 618]. Bu semantik mos kelmasliklar SI ijodidagi qisman so'ngan, xatoli kontentlar

shakllanishiga sabab bo'ladi [2008: 23]. Shu sababli, AI modellarida antonimiyani chuqurroq kognitiv va kontekstual ko'rinishda ta'limlash lozim. Bu uchun xalq ijodi, adabiy obrazlar, va metaforalar korpusini mashq jarayoniga kiritish taklif etiladi [2012: 143]. SI modellari tilni qayta ishlashida antonim juftliklarning embedding burchaklarini tahlil qilish foydali bo'ldi, lekin semiotik va ma'naviy yuklamalarni o'lchay olmaslik yetarlicha muammo sifatida qolmoqda [2016: 952]. Inson uchun "yoz – qish" jufti nafaqat ob-havo, balki turmush tarzi, madaniy o'zaro munosabat, holat ramzi hisoblanadi – AI modellarda bunday qabul yo'q [1986: 206]. Umumiy xulosa shuki, SI antonimiyani chuqur anglashi uchun u inson kognitiv faoliyatini, his-tuyg'uni va kontekstni qayd eta oladigan korpuslar bilan mashq qilishi kerak [1996: 167]. Bu jarayonda ko'p tilli, semiotik va milliy xususiyatlarga ega ma'lumotlar bazasi tayyorlanishi maqsadga muvofiq [2002: 78].

Ushbu tadqiqot asosida shunday xulosa qilish mumkinki, sun'iy intellekt modellari antonimiyani to'g'ri anglashda muayyan samaralarga erishgan bo'lsa-da, u hali inson tafakkuri kabi chuqur kognitiv va assotsiativ ma'nolarni qamrab olishda yetarli emas [2013: 581]. Formal leksik juftliklar yuqori aniqlikda tanilgan holda, metaforali va madaniy yuklamali juftliklarda samarasizlik kuzatildi [2012: 144]. Bunday holat SI modellarining ma'no tuzish jarayonida kontekst va emotsional bahoni hisobga olish qobiliyati zaif ekanligini ko'rsatadi [1996: 168]. Inson ma'noni jamiyat, tajriba va nutqiy vaziyat orqali anglar ekan, SI faqat korpusdagi statistik aloqalarga suyanmoqda [1986: 208].

Antonimiya inson ijodida ritorik vosita sifatida ham xizmat qilishi, ya'ni faqat ma'no emas, ta'sir va baholash funktsiyalarini ham bajarishi e'tiborga olinmoqda [1980: 125]. Shuning uchun AI modellarining antonimiyani to'g'ri anglash darajasi uning ijodkorlik qobiliyatini aniqlovchi mezonlardan biridir [2006: 400]. Antonimiyani modellashda faqat embedding vektorlarini taqqoslash kifoya emas – inson ongidagi semantik yo'llar va baholash konstruktleri ham hisobga olinishi lozim [2008: 24]. Bu uchun SI modellari adabiyot, xalq ijodi, badiiy obrazlar, va frazeologik birliklar asosida mashq qilinishi kerak [2002: 80].

Semantik qarshilikka ega so'z juftliklarini avtomatik anglash uchun ko'p tilli va kontekstual korpuslar yaratish – kelgusidagi asosiy vazifalardan biri hisoblanadi [2017: 618]. Masalan, "hayot–o'lim", "tinish – ko'z yoshi" kabi badiiy juftliklar faqat his-tuyg'u, tafakkur va ta'rifga bog'liq tarzda anglanadi [1981: 220]. Shu sababli, AI modellari uchun inson fe'l-atvorini aks ettiruvchi, kontekst sezuvchan semantik tuzilmalarni joriy etish talab etiladi [2000: 62]. Bu nafaqat antonimiyani anglash, balki umumiy tilni tushunish sifatini ham oshiradi [2016: 953]. Inson nutqidagi qarshilik munosabatlari turli darajada namoyon bo'ladi – fazoviy,

miqdoriy, baholovchi, psixologik, ijtimoiy, falsafiy va boshqa. SI modellari ushbu darajalarni tanlashda ko'pincha yuzaki tahlilga asoslanadi [2019: 9].

Antonimiya inson fikrlashining universal qonuniyatlaridan biri ekanligi hisobga olinsa, uni AI tizimlarida to'g'ri modellashtirish – sun'iy intellektning ma'naviy-kommunikativ jihatdan rivojlantirishning asosiy shartlaridan biriga aylanadi [1996: 170]. Shuning uchun, modellarni semantik intuitsiya bilan “ta’limlash”, ya’ni inson kognitiv tajribasini korpus orqali yuklash vazifasi dolzarb hisoblanadi [2020: 8]. Bu jarayonda semantik kontrastli korpuslar, emotsional leksika, va muloqot voqeligi aks etgan ma’lumotlar muhim o‘rin tutadi [2010: 620].

Ushbu maqolada taklif etilgan metodologik yondashuvlar – korpusli tahlil, ekspert baholash, embedding burchaklari tahlili – SIDA antonimiyani anglash darajasini baholashda foydalidir [2013: 582]. Shuningdek, natijalar inson va mashina o‘rtasidagi semantik uzilish nuqtalarini aniqlashga yordam beradi [2019: 10].

Adabiyotlar ro‘yxati:

1. Биск Й. и др. Experience Grounds Language // Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP). – 2020. – 12 с.
2. Веал Т. Exploding the Creativity Myth: The Computational Foundations of Linguistic Creativity. – London: Bloomsbury, 2012. – 256 с.
3. Девлин Дж., Чанг М.-В., Ли К., Тоутанова К. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding // NAACL-HLT. – 2019. – 11 с.
4. Джонс С. Antonymy: A Corpus-Based Perspective. – London: Routledge, 2002. – 240 с.
5. Лакофф Дж., Джонсон М. Metaphors We Live By. – Chicago: University of Chicago Press, 1980. – 256 с.
6. Лиу Й. и др. RoBERTa: A Robustly Optimized BERT Pretraining Approach // arXiv preprint arXiv:1907.11692. – 2019. – 15 с.
7. Макрей Дж., Бюителар П. Expanding wordnets using unsupervised methods // Language Resources and Evaluation. – 2017. – Т. 51. – №3. – С. 603–627.
8. Мерфи М. Semantic Relations and the Lexicon: Antonymy, Synonymy and Other Paradigms. – Cambridge: Cambridge University Press, 2003. – 296 с.
9. Мохаммад С., Дорр Б., Хёрст Г., Терни П. Computing Lexical Contrast // Computational Linguistics. – 2013. – Т. 39. – №3. – С. 555–590.

10. Петерс М. и др. Deep Contextualized Word Representations // NAACL. – 2018. – 9 с.
11. Рикоёр П. The Rule of Metaphor: The Creation of Meaning in Language. – London: Routledge, 1981. – 384 с.
12. Райзингер Дж., Муни Р. Multi-Prototype Vector-Space Models of Word Meaning // NAACL. – 2010. – 8 с.
13. Росч Э. Principles of Categorization // In: Cognition and Categorization. – Hillsdale: Lawrence Erlbaum, 1978. – С. 27–48. – 21 с.
14. Счультэ им Вальде С., Кисселев М. Distributional Approaches to Semantic Relatedness // Language Resources and Evaluation. – 2016. – Т. 50. – №4. – С. 935–963.
15. Талми Л. Toward a Cognitive Semantics. – Cambridge: MIT Press, 2000. – Т. 1. – 563 с.
16. Терни П. Similarity of Semantic Relations // Computational Linguistics. – 2006. – Т. 32. – №3. – С. 379–416.
17. Уирзбицка А. Semantics: Primes and Universals. – Oxford: Oxford University Press, 1996. – 512 с.
18. Флауэрс Е., Голдберг Й. Adversarial Examples for Evaluating Reading Comprehension Systems // TACL. – 2018. – Т. 6. – С. 605–617.
19. Чатжи Ю., Лиу Й. и др. RoBERTa: A Robustly Optimized BERT Pretraining Approach // arXiv preprint. – 2019. – 15 с.
20. Шультэ им Вальде С., Кисселев М. Distributional Approaches to Semantic Relatedness // Language Resources and Evaluation. – 2016. – Т. 50. – №4. – С. 935–963.
21. Элэзар Й., Голдберг Й. Adversarial Examples for Evaluating Reading Comprehension Systems // Transactions of the ACL. – 2018. – Т. 6. – С. 605–617.
22. Элазер Й., Голдберг Й. Adversarial Examples for Evaluating Reading Comprehension Systems. – 2018. – 13 с.
23. Cruse D. A. Lexical Semantics. – Cambridge: Cambridge University Press, 1986. – 328 p.
24. Turney P. Similarity of semantic relations // Computational Linguistics. – 2006. – Vol. 32. – No. 3. – P. 379–416.
25. Mohammad S. et al. Obtaining reliable human ratings of valence, arousal, and dominance for 20,000 English words // LREC. – 2018. – 7 p.