

Ayxan Ədalət oğlu Mustafayev

AMEA Fəlsəfə və Sosiologiya İnstitutunun doktorantı

E-mail: ayxan.mustafayev20@gmail.com

ORCID ID: [0009-0000-8536-9883](https://orcid.org/0009-0000-8536-9883)

Elmi rəhbərlər:

fəlsəfə elmləri doktoru, baş elmi işçi Füzuli Məhəmməd oğlu Qurbanov

fəlsəfə üzrə fəlsəfə doktoru, dosent Məhəmməd Sultan oğlu Cəbrayilov

BÖYÜK DİL MODELƏRİ VƏ LİŦQVİSTİK BİLİK PROBLEMİ: ANLAMA, YOXSƏ SİMULYASIYA

Məqalədə böyük dil modellərinin (LLM) linqvistik bilik problemi epistemoloji və fəlsəfi müstəvidə araşdırılır. Tədqiqatın əsas məqsədi ChatGPT, Claude, Gemini və digər böyük dil modellərinin insan dilini istehsal etmə qabiliyyətinin həqiqi anlama, yoxsa statistik simulyasiya ilə bağlı olduğunu müəyyənləşdirmək, eyni zamanda bu sualın həllinin biliyin və mənanın təbiəti barədə klassik təsəvvürlərə necə təsir etdiyini göstərməkdir. Məqalədə əsaslandırılır ki, bu sistemlərin qrammatik cəhətdən koherent və məzmunca mənalı görünən mətnlər yaratması hələ onların dili insan kimi anladığını sübut etmir. Sintaktik uğurla semantik qavrayış arasındakı fərq diqqətdən kənarda qalmamalıdır. Bununla yanaşı, LLM-ləri yalnız “statistik tutuquşu” kimi izah etmək də yetərli deyil, çünki bu yanaşma fenomenin epistemoloji mürəkkəbliyini və onun idrak proseslərinə real təsirini tam əhatə etmir. Məqalədə Wittgenstein, Searle, Chomsky, Floridi, Bender və digər tədqiqatçıların yanaşmaları tənqidi şəkildə müqayisə edilərək, problemi iki uc mövqe olan tam anlama və sırf simulyasiya arasında üçüncü mövqedən izah edilir və bu “paylanmış linqvistik idrak” konsepsiyası kimi təqdim edilir. Bu konsepsiyaya görə, məna nə yalnız insan zehmində, nə də yalnız alqoritmik modelin daxilində yerləşir. O, insanla maşın arasındakı interaktiv, dialoji münasibətdə və istifadə kontekstində formalaşır. Məqalədə həmçinin LLM-lərin halüsinasiya problemi, biliyin doğruluğuna görə epistemik məsuliyyətin kimə aid olması, müəlliflik və orijinallıq məsələləri, eləcə də Azərbaycan dili kimi az resurslu dillər kontekstində bu modellərin yaratdığı imkanlar və risklər təhlil olunur. Nəticə olaraq qeyd edilir ki, böyük dil modelləri biliklər cəmiyyətində dilin, mənanın və biliyin statusunu, habelə insanın idrak prosesindəki mövqeyini yenidən düşünməyi zəruri edir.

Açar sözlər: böyük dil modelləri, Linqvistik bilik, Süni intellekt, Simulyasiya, Epistemologiya, Paylanmış linqvistik idrak, Biliklər cəmiyyəti.

Giriş

Son illərdə süni intellekt sahəsində ən mühüm inkişaflardan biri böyük dil modellərinin sürətli inkişafı və gündəlik həyatın müxtəlif sahələrinə geniş inteqrasiyası olmuşdur. ChatGPT, Claude, Gemini, DeepSeek və digər sistemlər artıq yalnız texniki alətlər kimi deyil, həm də mətn yaratma, tərcümə, proqramlaşdırma, təhlil, bilik istehsalı və hətta fəlsəfi müzakirələrdə iştirak edən yeni rəqəmsal vasitələr kimi fəaliyyət göstərir. Bu sistemlər sualları cavablandırır, mürəkkəb mətnlər qurur, arqumentlər formalaşdırır və insan dilinə yaxın ifadə imkanları nümayiş etdirir. Nəticədə onların texniki imkanlarının genişlənməsi klassik fəlsəfi sualı yenidən aktuallaşdırır: böyük dil modelləri həqiqətən dili başa düşür, yoxsa sadəcə anlama illüziyası yaradan statistik simulyasiya mexanizmləridir?

İlk baxışda bu sual texnologiyə problem kimi görünə bilər. Lakin onun əsasında daha dərin epistemoloji və ontoloji məna dayanır. Çünki “anlama” anlayışı fəlsəfə tarixində heç vaxt sadə və yekdil qəbul edilən kateqoriya olmayıb. Dil və idrak, məna və təcrübə, bilik və ifadə arasındakı əlaqə Platondan Vitgenşteynə, Çomskidən Searleyə qədər müxtəlif nəzəri yanaşmalarda müzakirə olunmuşdur. Müasir dövrdə isə bu məsələ insan və maşın arasındakı sərhədin daha qeyri-müəyyən hala gəldiyi yeni şəraitdə ortaya çıxır. Buna görə böyük dil modellərinin mahiyyəti yalnız texnologiyaya fəlsəfəsinin deyil, ümumilikdə bilik fəlsəfəsinin də fundamental problemlərindən biri kimi araşdırılmalıdır.

Məsələnin aktuallığı dilin insan idrakı və sosial bilik istehsalı üçün əsas vasitələrdən biri olması ilə bağlıdır. İnsan biliyi yalnız fərdi şüurda mövcud olmur; o, dil vasitəsilə paylaşılır, ötürülür, əsaslandırılır və institusional məna qazanır. Dil yarada bilən qeyri-insani sistemlərin meydana çıxması biliklərin formalaşması, yayılması və legitimləşdirilməsi mexanizmlərinə təsirsiz qala bilməz. Bu səbəbdən böyük dil modellərinin “başla düşüb-düşməməsi” sualını sadə bəli və ya yox cavabı ilə həll etmək kifayət deyil. Belə yanaşma həm bu sistemlərin məhdudiyyətlərini, həm də bilik cəmiyyətində yaratdıqları yeni epistemik vəziyyəti tam əhatə etmir.

Bu məqalədə əsas argument ondan ibarətdir ki, böyük dil modelləri insanlarla eyni şəkildə anlama qabiliyyətinə malik deyil, lakin onları sadəcə mexaniki təqlid və ya statistik simulyasiya kimi təqdim etmək də məsələni həddindən artıq sadələşdirir. Onların yaratdığı mətnə məna nə tamamilə modelin daxilində, nə də yalnız istifadəçinin şüurunda yerləşir; məna insan və maşın arasındakı qarşılıqlı təsir prosesində formalaşır. Buna görə “anlama, yoxsa simulyasiya?” sualı kəskin dioxotomiya kimi deyil, insan–maşın münasibətləri daxilində yeni linqvistik və epistemik təcrübələrin yaranması baxımından təhlil edilməlidir. Əgər müasir bilik yalnız fərdi subyektin məhsulu deyil, həm də insanların, alqoritmlərin və platforma infrastrukturalarının birgə təsiri ilə formalaşırsa, bu prosesin əsas ifadə vasitəsi olan dilin statusu da yenidən araşdırılmalıdır.

Bu kontekstdə məqalənin məqsədi böyük dil modellərində dil bilikləri problemini epistemoloji və fəlsəfi baxımdan təhlil etməkdir. Tədqiqat əvvəlcə dil və bilik arasındakı əlaqəni klassik epistemoloji çərçivədə araşdırır, daha sonra böyük dil modellərinin texniki və fəlsəfi mahiyyətini izah edir və “anlama” ilə “simulyasiya” mövqelərinin müqayisəli təhlilini aparır. Növbəti mərhələdə bu mübahisənin izahı üçün “paylanmış dil idrakı” anlayışı təklif olunur. Sonda isə bu yanaşmanın Azərbaycanda bilik cəmiyyəti, epistemik məsuliyyət, müəlliflik, etika və akademik mühit üçün nəticələri müzakirə edilir. Məqalənin metodoloji əsasını tənqidi nəzəri təhlil təşkil edir və araşdırma həm klassik fəlsəfi mətnlərə, həm də süni intellekt və dil modelləri ilə bağlı müasir tədqiqatlara əsaslanır. Bu yanaşma böyük dil modellərini yalnız texniki sistemlər kimi deyil, bilik, məna, dil və məsuliyyət münasibətlərini yenidən formalaşdıran fəlsəfi problem kimi başla düşməyə imkan verir.

1. Dil və Bilik: Klassik Epistemoloji Çərçivə

Dilin epistemoloji statusu: Fəlsəfə tarixi boyunca dil heç vaxt sadəcə ünsiyyət vasitəsi kimi qəbul edilməyib. Platondan başlayaraq dil də bilik strukturuna təsir edən bir mexanizm kimi qəbul edilib. Aristotelin Organondakı qrammatik təhlilləri, orta əsr İslam filosoflarının (Əl-Fərabî, İbn Sina) mənanın təbiəti haqqında əsərləri və daha sonra Avropada Anglo-Sakson analitik fəlsəfəsi dilə xüsusi epistemoloji əhəmiyyət vermişdir. Bütün bu ənənələrdə ortaq bir xüsusiyyət fərqlənir: dil həm bəşəriyyətin dünya ilə münasibətini, həm də bu əlaqənin bilik şəklində ifadəsini şərtləndirir.

Bu ənənədə xüsusi yer tutan filosoflardan biri də Lüdviq Vitgenşteyndir. O, “Tractatus Logico-Philosophicus” əsərində məşhur tezisini irəli sürmüşdür: “Dilimdəki məhdudiyyətlər dünyamdakı məhdudiyyətlər deməkdir” [20, təklif 5.6]. Bu fikir dilin yalnız dünyanı təsvir etmədiyini, həm də onu qurduğunu göstərir. Sonrakı “Fəlsəfi Tədqiqatlar” əsərində

Vitgenşteyn bu mövqeyi daha da inkişaf etdirərək mənanın istifadədə yarandığını iddia etmişdir: "Sözün mənası onun dildə istifadəsidir" [20, s. 43]. Bu yanaşma, dil istifadəsinin böyük statistik nümunələrinə əsaslanan müasir linqvistik öyrənmə modellərinin (LLM) müzakirələri üçün xüsusilə əhəmiyyətlidir. Daha da maraqlısı odur ki, əgər məna istifadədə yaranırsa, LLM-lərin bu istifadəyə təsiri epistemoloji nəticələrə malik olmalıdır.

Maraqlıdır ki, mənanın dildə təşəkkülü məsələsi Azərbaycan epistemoloji ədəbiyyatında da müzakirə olunmuşdur. Məmmədov və Qurbanov Lütfi Zadənin nəzəriyyəsini təhlil edərək belə nəticəyə gəlirlər ki, "düşüncə prosesində insan öz fikrini demək üçün əslində onu qeyri-səlis çoxluqların köməyi ilə approksimasiya edir" [3, s. 27]. Bu yanaşma mənanın insan təfəkküründə qəti, ikili (binar) dəyərlərlə deyil, dərəcəli və qeyri-səlis şəkildə formalaşdığını göstərir. Bu fikir mənim sonradan irəli sürəcəyim "anlamanın dərəcə məsələsi olması" argumenti ilə birbaşa səsleşir.

Noam Çomski isə fərqli bir yanaşma tətbiq etmişdir. O, Dilin fitri quruluşu və generativ qrammatika nəzəriyyəsini inkişaf etdirərək, insan dilinin yalnız statistik öyrənmə yolu ilə deyil, spesifik idrak memarlığı tələb etdiyini nümayiş etdirmişdir [9]. Çomskinin fikrincə, dil istifadəsi və daxili qrammatik bilik iki fərqli şeydir: "birincisi performans, ikincisi isə səriştədir" [10]. Bu fərq mövzumuzla birbaşa əlaqəlidir. Beləki, dil öyrənenlər "dil performansı" nümayiş etdirirlər, lakin onların həqiqətən "dil bacarıqları" qazandıqlarını demək çətindir. Çomskinin buradakı fikri ilə tamamilə razı olmasam da, argumentimdə performans-səriştə fərqlərini qorumaq lazım olduğunu düşünürəm.

Klassik epistemologiya çərçivəsində dil daha bir vacib funksiya yerinə yetirir. O, biliyin sosial ötürülmə vasitəsidir. Şəxsi təcrübə fərdiyyətlə məhdudlaşır, lakin dildə ifadə olunmuş bilik nəsillər arasındakı ötürülmə bilirdir. Karl Popper bunu "üçüncü dünya", yəni, obyektiv biliyin yaşadığı sahə adlandırmışdır [17, s. 106]. Bu mövqedən baxılırsa, LLM-lər "üçüncü dünyaya" yeni bir aktyor kimi qatılır. Beləki, onlar dildə ifadə olunmuş bilik kütlələrini mənimsəyir və onları yeni formada təqdim edir. Lakin sual qalır: bu yeni təqdimat həqiqi mənada bilik istehsalıdır, yoxsa yalnız onun statistik proyeksiyasıdır? Bu suala cavab axtarmaq həmin proyeksiyanın daxili strukturunu təhlil etməyi tələb edir.

2. Böyük Dil Modelləri: Texniki və Fəlsəfi Səciyyə

LLM-lərin işləmə prinsipi və epistemoloji nəticələri: Böyük dil modelləri nədir? Texniki cəhətdən, onlar adətən "transformator" arxitekturasına əsaslanan milyardlarla parametrlə malik nəhəng neyron şəbəkələridir, yəni statistik sistemlərdir. Onların əsas funksiyası olduqca sadədir: müəyyən bir kontekstdə növbəti sözün (tokenin) nə olması sualına cavab vermək. Murray Shanahan bu prosesi belə izah edir: "Böyük dil modelləri mənanı deyil, statistik nümunələri öyrənir; yaratdıqları mətn dilin formasına bənzəyir, lakin bu, dilin özü deyil" [19, s. 71]. Bu izah məsələni həddindən artıq sadələşdirsə də, əsas məqamı əhatə edir: böyük dil modellərinin daxili işi mənalara deyil, ehtimallara əsaslanır.

Lakin, texniki sadəliklərinə baxmayaraq, LLM-lərin davranışı təəccüblü dərəcədə mürəkkəbdir. Tədqiqatçılar "ortaya çıxan qabiliyyətlər" (emergent capabilities) kimi tanınan bir fenomen müşahidə ediblər: model böyüdükcə, ona xüsusi olaraq öyrədilməyən bacarıqlar (məsələn, tərcümə, məntiqi mühakimə və kodlaşdırma) qəfildən ortaya çıxır [8, s. 8]. Bu, statistik modelin "anlamaya" doğru irəlilədiyini göstərə bilsə də, fenomenin daha dərin təhlili göstərir ki, sözdə "ortaya çıxan" davranışlar hələ də təlim məlumatlarındakı nümunələrlə izah olunur. Başqa sözlə, biz "sehrli" bir fenomenlə deyil, yeni xüsusiyyətlər yaradan statistik miqyaslaşma ilə məşğuluq.

Burada fəlsəfi baxımdan maraqlı bir sual ortaya çıxır: kəmiyyət keyfiyyətə çevrilirmi? Hegelin dialektikasında bu, klassik bir prinsipdir - kəmiyyətdəki dəyişikliklər müəyyən bir nöqtədə keyfiyyət sıçrayışına səbəb olur. LLM-lərdə müşahidə olunan ortaya çıxan imkanlar bu prinsipin texnoloji təzahürü kimi şərh edilə bilər. Lakin ehtiyatlı olmaq lazımdır: bu

prinsipin LLM-lərə tətbiqi sadəcə bir metaforadır, çünki kompüter sistemindəki "keyfiyyət sıçrayışı" fəlsəfi mənada "keyfiyyət" deyil, sadəcə davranış nümunələrinin yeni bir təşkilatıdır. İnanıram ki, bu nüans tez-tez nəzərdən qaçırılır və LLM-lərə lazımsız fəlsəfi əhəmiyyət verilməsinə gətirib çıxarır.

Dilçilik öyrənmə modellərinin necə işlədiyinin digər əsas aspekti kontekst pəncərəsinin məhdudiyyətidir. Model hər sorğuda yalnız müəyyən sayda işarəni emal edə bilər; bundan kənarda olan məlumatlar "unudulur". Bu məhdudiyyət onları insan dilinə bənzər proseslərdən fərqləndirən ən fərqli xüsusiyyətdir, çünki insan dil təcrübəsi bütün həyat hekayəsini, mədəni yaddaşı və davamlı kimliyi əhatə edir. Lakin dil öyrənmə modellərində belə uzunmüddətli ardıcılıq yoxdur; hər sessiya yeni bir başlanğıcdır. Bu, "anlamanın" baş verdiyi şərtləri müzakirə edərkən mütləq nəzərə alınmalı vacib bir məqamdır.

"Stochastic Parrot" tezisi və onun məhdudiyyətləri: Emily Bender, Timnit Gebru və həmkarları tərəfindən 2021-ci ildə təqdim edilən "stochastic parrots" (statistik tutuquşular adlı məqalə böyük dil modelləri (LLM) ətrafındakı müzakirələrdə dönüş nöqtəsi oldu. Müəlliflərə görə: "LLM-lər linqvistik formanı manipulyasiya etmək üçün böyük verilənlər bazalarından istifadə edən sistemlərdir; onlar mənaya istinad etmirlər, çünki onların mənə ilə heç bir əlaqəsi yoxdur" [6 s. 616–617]. Bu baxışa görə, LLM-lər tərəfindən yaradılan mətnin uyğunluğu tamamilə dil istifadəçilərinin, yəni biz insanların həmin mətnə aid etdiyi mənadan asılıdır.

Bender və Koller başqa bir məqalədə daha güclü bir arqument irəli sürmüşdülər: dil sadəcə forma məsələsi deyil; mənə onun ayrılmaz hissəsidir; mənə dilin dünyadakı təməlinə (grounding) yaranır. Dil öyrənmə maşınları (LLM) yalnız formaya daxil ola bilər, dünyaya yox; buna görə də onlar təbii olaraq mənəni ifadə edə bilməzlər [7, s. 5188]. Bu arqument güclüdür və mənə onun əsas məntiqini qəbul edirəm. Lakin, tezisə bir məhdudiyyəti var: o, "mənə"nin mütləq fiziki təməllə bağlı olduğunu fərz edir. Lakin insanlar konkret fiziki təməl olmadan mücərrəd anlayışları (məsələn, "ədalət" və ya "sonsuzluq") başa düşürlər. Buna görə də, təməl anlayış üçün yeganə şərt deyil.

Benderin tezisindəki digər bir məhdudiyyət "tutuquşu" metaforasının təbiətindədir. Tutuquşu sözləri təkrarlayır, lakin LLM-lər sadəcə təkrar etmir; onlar yeni kombinasiyalar yaradır və əvvəllər heç vaxt qarşılaşmadıqları suallara yeni cavablar yaradırlar. Bu xüsusiyyət "statistik tutuquşu" metaforasının altında gizlənilir və məsələni həddindən artıq sadələşdirir. Burada Benderin haqlı olduğunu qəbul edirəm - LLM-lər insanlar kimi başa düşürlər - amma onun metaforası fenomenin yaradıcı tərəfini görməyə imkan vermir. Bu məhdudiyyət daha sonra təklif edəcəyim yanaşma üçün vacib bir başlanğıc nöqtəsidir.

"Halüsinasiyalar" problemi: "Halüsinasiyalar" problemi böyük dil modellərinin epistemoloji mahiyyətini aydın göstərən əsas xüsusiyyətlərdən biridir. Bu modellər bəzən tamamilə yanlış və ya uydurma məlumatları inamlı şəkildə təqdim edir, çünki onlar həqiqət və yalanı daxili kateqoriyalar kimi ayırd etmir, sadəcə sözlərin ardıcılığını statistik ehtimallar əsasında qururlar. Klassik epistemologiyada bilik "doğru və əsaslandırılmış inanc" kimi başa düşüldüyü halda, LLM-lərdə nə inanc, nə də doğruluq meyarı mövcuddur. Buna görə onlar dar mənada bilik istehsal etmir, lakin biliyə bənzər mətnlər yaratdıqları üçün istifadəçilər tərəfindən çox vaxt bilik kimi qəbul olunurlar. Məhz bu uyğunsuzluq biliklər cəmiyyətində ciddi risk yaradır: yanlış məlumat doğru məlumat forması olaraq dövryyəyə daxil olur və bu, LLM-lərin həqiqi mənada "anlamadığını" göstərən ən güclü praktik dəlillərdən biridir.

3. Anlama, yoxsa Simulyasiya: Mərkəzi Mübahisə

Searle-in "Çin otağı" arqumenti və müasir tətbiqi: Bu müzakirənin fəlsəfi əsasını anlamaq üçün Con Searlin ilk dəfə 1980-ci ildə təklif etdiyi məşhur "Çin otağı" düşüncə təcrübəsinə qayıtmalıyıq. Searl bizdən Çin dilini bilməyən bir insanın otaqda oturub, simvollar yaratmaq və cavab göndərmək üçün müəyyən qaydalara uyğun olaraq Çin hərflərini

manipulyasiya etdiyini təsəvvür etməyimizi xahiş edir. Kənar müşahidəçiyə bu şəxs Çin dilini bilirmiş kimi görünür; lakin əslində onlar simvolların mənasından xəbərsizdirlər. Searl yazır: "Simvolların manipulyasiyası sintaksisdir, lakin sintaksis özlüyündə semantika təmin etmir" [18, s. 423].

Bu arqument bütün LLM-lərə aiddir. Modellər ehtimallara əsaslanaraq növbəti işarələri seçərək söz vektorlarını manipulyasiya edirlər; lakin, bu prosesdə "məna" kimi bir şey yoxdur. Searle-nin "sintaksis məna yaratmır" tezi, LLM-lərin "başə düşmədiyi" nəticəsini dəstəkləyən ən güclü arqumentlərdən biridir. Lakin, mən ümumiyyətlə Searle-nin arqumentini qəbul etsəm də, onun son nəticəsi ilə tam razı deyiləm. Razı olmadığım məqam, Searle-nin prinsipə, hər hansı bir qeyri-insani sistemin başə düşə biləcəyi ehtimalını rədd etməsidir. Lakin məsələ daha incədir: başə düşməyin mümkün olub-olmaması yalnız sistemin daxili xüsusiyyətlərindən deyil, həm də digər sistemlərlə qarşılıqlı təsirindən asılıdır.

Searle-nin arqumentinə qarşı ən məşhur etiraz "sistem reaksiyası"dır: otaqdakı şəxs Çin dilini bilmir, lakin bütün "otaq+şəxs+qaydalar" sistemi Çin dilini "bilir". Searle bu etirazı qətiyyətlə rədd edərək, sistemə nə əlavə etməyinizdən asılı olmayaraq, sintaksisdən məna çıxarmağın mümkün olmadığını iddia etdi. Bu arqumentdə mən Searle-ə daha yaxınam, amma yeni bir nüansı vurğulayıram: burada əsas məqam LLM-ə "müstəqil sistem" kimi deyil, "şəxs + LLM" cütünü kimi yanaşmaqdır. Bu cütlükdə məna şəxs tərəfindən təmin edilir, lakin onun ifadəsi yeni hibrid formada baş verir. Bu nüans Searle-in arqumentini zəiflətmir, əksinə onu gücləndirir.

"Anlama" lehinə arqumentlər: Funksionalistlər, Searle-ə qarşı çıxaraq, bir sistemin davranışını funksional cəhətdən mənalı bir sistemin davranışından ayırd etmək mümkün olmadığı təqdirdə, sistemin də şüurlu olduğunu düşünməliyik. Başqa sözlə, şüur daxili bir xüsusiyyət deyil, müşahidə edilə bilən bir davranışdır. Bu baxış Daniel Dennett və başqaları tərəfindən müdafiə edilmişdir [13, s. 89]. Funksionalizmin gücü onun davranış qiymətləndirməsinə əsaslanmasındadır: əgər heyvanlara "anlama" aid etsək - hətta onların daxili təcrübələrini bilə bilməsək də - onda niyə uzunmüddətli məntiqi sistemlər üçün eyni şeyi etməməliyik?

Funksionalizmə qarşı klassik etiraz, davranış oxşarlığının eynilik demək olmamasıdır. İki sistem eyni davranış modelini nümayiş etdirə bilər, lakin tamamilə fərqli daxili prosesləri idarə edə bilər. Bu, "zombi arqumenti" kimi tanınır: insan kimi davranan, lakin daxili təcrübəsi olmayan bir varlıq təsəvvür edin. Funksionalistlərə görə, bu varlıq davranış eyni olduğu üçün "başə düşür". Lakin intuisiyamız bunu qəbul etməyə qəti şəkildə müqavimət göstərir. Mənim fikrimcə, linqvistik media məhz belə bir "zombi linqvistik agent"dir; insanın linqvistik davranışını təqlid edən, lakin daxili təcrübəsi olmayan bir agent.

Luciano Floridi daha incə bir yanaşma təklif edir. "Zehin fəlsəfəsi" çərçivəsində o, süni intellekt sistemlərinin yeni bir fəaliyyət forması təşkil etdiyini iddia edir: şüurlu olmasalar belə, mənalı şəkildə fəaliyyət göstərə bilən, lakin insanlardan fərqli şəkildə fəaliyyət göstərən sistemlər. Floridi yazır: "Süni intellekt zehni xüsusiyyətlərdən məhrum olan yeni bir fəaliyyət formasıdır; o, fəaliyyət göstərə bilər, lakin anlayış ona bizimlə eyni şəkildə tətbiq olunmur" [14, s. 12]. Bu yanaşma Searle-nin arqumentini tam qəbul etmir, çünki "anlama"nı yalnız insan modeli baxımından təyin etməyə etiraz edir. Floridi-nin yanaşması mənim üçün daha çox rezonans doğurur, çünki o, "ya ya/ya da" məntiqindən uzaqlaşır. Paylama semantik yanaşması da diqqətəlayiqdir. Bu yanaşmaya görə, bir sözün mənası digər sözlərlə yanaşı istifadə nümunələrindən yaranır. LLM-lər bu prinsip üzərində qurulub və əgər məna paylanmış nümunələrdən ortaya çıxarsa, onda LLM-lərin müəyyən mənada məna ilə fəaliyyət göstərdiyi iddia edilə bilər [16, s. 11]. Bu yanaşma Vitgenşteynin "məna istifadədə ortaya çıxır" tezi ilə uyğun gəlir, lakin onu radikal ifrat səviyyəyə qaldırır. Bununla belə, ehtiyatlı

olmaq vacibdir: Vitgenşteyn dil istifadəsini insan həyat tərzinin bir hissəsi kimi görürdü; o, dili insanlardan ayrı bir kateqoriyaya endirməmişdir.

Paylanmış Linqvistik İdrak: Paylanmış Linqvistik İdrak: Hər iki tərəfin arqumentlərini nəzərə aldıqda görünür ki, həm Benderin “təsadüfi tutuquşular” mövqeyində, həm də Floridinin daha çevik “anlama” yanaşmasında müəyyən həqiqət payı vardır, lakin onların heç biri problemi tam izah etmir. Bender haqlıdır ki, LLM-lər insanlarla eyni mənada başa düşmür; Floridi isə haqlı olaraq göstərir ki, anlamayı yalnız insan modeli ilə məhdudlaşdırmaq fəlsəfi baxımdan yetərli deyil. Bu iki mövqe arasında üçüncü yol kimi “paylanmış linqvistik idrak” anlayışını təklif edirəm. Bu yanaşmaya görə, böyük dil modellərində məna nə yalnız modelin daxilində, nə də yalnız insan istifadəçisinin zehninə yerləşir; məna insan–maşın qarşılıqlı təsiri prosesində yaranır. Buna görə “Böyük dil modelləri dili başa düşürmü?” sualından daha düzgün sual budur: “İnsan–LLM hibridi yeni linqvistik təcrübə yaradırmı?” Bu konsepsiya Clark və Chalmers-in “genişlənmiş zəhn” və Hutchins-in “paylanmış idrak” yanaşmalarından ilham alsa da, onlardan fərqlənir. Çünki LLM burada sadəcə köməkçi alət deyil: o, aktiv şəkildə dil yaradır, lakin niyyətə malik olmadığı üçün tərəfdaş da sayıla bilməz. Bu baxımdan LLM-ləri “aktiv linqvistik mühit” kimi başa düşmək daha uyğundur. Beləliklə, epistemik agentlik insan–LLM cütünü arasında paylanır və bu paylaşımın əsas təzahürü dil istehsalı prosesində üzə çıxır.

Maraqlıdır ki, insan–maşın münasibətinin bu qarşılıqlı təbiəti yerli tədqiqatlarda da diqqət çəkmişdir. Mahmudov (2024) yaygın “təqlid” fərziyyəsini tərsinə çevirərək göstərir ki, “süni intellekt insan düşüncəsini təqlid etmir, əksinə insan süni intellekt səviyyəsində müəyyən əməliyyatlar aparmağa cəhd göstərir” [2, s. 11]. Başqa sözlə, dilçi və ya istifadəçi öz təfəkkür tərzini sistemin emal edə biləcəyi formal kateqoriyalara uyğunlaşdırır. Bu müşahidə paylanmış linqvistik idrak konsepsiyasını gözlənilməz bir istiqamətdən təsdiqləyir: əgər təkcə maşın insana deyil, insan da maşına uyğunlaşarsa, onda dil praktikasını həqiqətən tək tərəfli deyil, qarşılıqlı və hibrid xarakter daşıyır.

Bu yanaşmanın üstünlüyü ondadır ki, o, həm Searle-in “sintaks semantika yaratmaz” tezisinə qarşı çıxmır, həm də LLM-lərin əhəmiyyətini azaltmır. Belə ki, LLM özlüyündə dili anlamır, lakin insan–LLM cütünü dil praktikasının yeni formasını yaradır. Mark Coeckelbergh-in dediyi kimi: “Süni intellektin epistemik təsiri onun nə olduğundan deyil, onun insanlarla necə qarşılıqlı əlaqədə olduğundan asılıdır” [12, s. 1349]. Mənim konsepsiyam bu fikri linqvistik kontekstə tətbiq edir və onu bir az da inkişaf etdirir: dil praktikasını özü artıq sırf insan praktikasını olaraq qalmır, o, hibrid xarakter qazanır.

Bu yanaşmanın digər bir fəlsəfi nəticəsi yeni bir növ “anlama”nın mümkünlüyüdür. İnsan anlayışı bir növ, heyvan anlayışı başqa bir növ, LLM anlayışı isə üçüncü növ “anlama”dır. Hər üçü özlüyündə etibarlıdır, lakin onları eyni ölçü ilə qiymətləndirmək səhvdir. Bu, fəlsəfi plüralizmə doğru bir addımdır. Vurğulanmalı olan məqam budur: plüralizm “hər şey etibarlıdır” mənasında nisbilik deyil; biz yenə də doğru ilə yanlış, müəyyən ilə qeyri-müəyyən anlayış arasında fərq qoymalıyıq. Lakin bu fərqi qoymaq üçün meyarlar yenidən araşdırılmalıdır.

4. Biliklər Cəmiyyətində Dilin Yeni Statusu:

Epistemoloji nəticələr: Epistemoloji nəticələr baxımından paylanmış linqvistik idrak konsepsiyası informasiya cəmiyyətində dilin statusunu yenidən düşünməyi zəruri edir. Böyük dil modellərinin ortaya çıxması göstərir ki, dil artıq yalnız insanın fərdi ifadə vasitəsi kimi deyil, həm də texnoloji sistemlər, verilənlər bazaları, alqoritmik emal mexanizmləri və istifadəçi müdaxilələri arasında formalaşan mürəkkəb idrak sahəsi kimi nəzərdən keçirilməlidir. Bu isə klassik epistemoloji kateqoriyaların, xüsusilə bilik, müəlliflik, subyekt və məsuliyyət anlayışlarının yenidən qiymətləndirilməsini tələb edir. Ən mühüm dəyişikliklərdən biri bilik istehsalında “müəllif” anlayışının transformasiyasıdır. Klassik

epistemologiyada bilik adətən konkret bir subyektə, yəni düşünən və yaradan şəxsə aid edilirdi. Lakin linqvistik öyrənmə modelləri tərəfindən yaradılan mətnlərdə müəlliflik artıq tək bir şəxsin fəaliyyəti kimi izah oluna bilmir. Belə mətnlərdə istifadəçinin sualı, modelin alqoritmik strukturu, təlim verilənlərinin məzmunu və həmin verilənlərə töhfə vermiş milyonlarla anonim mənbə birlikdə iştirak edir. Buna görə də sual yaranır: belə bir mətnin müəllifi kimdir? İstifadəçi, model, yoxsa modeli formalaşdıran geniş və paylanmış informasiya mühiti? Bu sual ilk baxışda texniki məsələ kimi görünə bilər, lakin onun dərin sosial, akademik və hüquqi nəticələri vardır. Universitetlərdə plagiat anlayışı, elmi nəşrlərdə müəlliflik etikası, tədqiqat məsuliyyəti və hüquqi mənada müəlliflik hüququ bu yeni şəraitdə əvvəlki kimi izah oluna bilmir. Böyük dil modelləri ilə yaradılan mətnlər klassik fərdi müəllif modelinin sərhədlərini zəiflədir və bizi paylanmış müəlliflik modelini nəzərə almağa məcbur edir. Bu keçid qaçılmaz görünsə də, onun hansı normativ və institusional mexanizmlərlə tənzimlənməyəyi hələ açıq məsələ olaraq qalır. Xüsusilə yerli akademik mühitdə süni intellektlə yaradılan mətnlərin müəllifliyi, etik istifadəsi və məsuliyyət bölgüsü ilə bağlı aydın çərçivəyə ehtiyac vardır.

Etibarlılıq və epistemik məsuliyyət: Digər vacib tapıntı informasiyanın etibarlılığı ilə bağlıdır. Daha əvvəl qeyd etdiyim böyük dil modellərindəki (LLM) "hallüsinasiya" problemi göstərir ki, bu sistemlər doğru ilə yanlış arasındakı fərqi bilmir. Nəticə etibarilə, informasiya cəmiyyəti daxilində bu sistemlərə etibar etmək yeni bir epistemik risk yaradır. Bu risk təkcə texniki deyil, həm də fəlsəfi və institusionaldır. Doktorluq dissertasiyamda göstərdiyim kimi, epistemik səriştə insan tədqiqatçıları, alqoritmik sistemlər və institusional infrastrukturлар arasında bölüşdürülərsə, məsuliyyət də bölüşdürülməlidir. Lakin, praktikada "paylanmış məsuliyyət" çox vaxt "heç kimin məsuliyyəti"nə bərabər ola bilər. Bu məsələ müasir süni intellekt etikasında "hesabat boşluğu" kimi adlandırılır və hələ də həll olunmayıb. Mənim təklifim budur ki, LLM-lərin istifadəsi üçün məsuliyyət son istifadəçidə qalmalı, sistemləri yaradan və yayan şəxslər də mütənasib məsuliyyət daşımalıdırlar. Bu, hüquqi və etik normaların yenidən qurulmasını tələb edir.

Azərbaycan kontekstində perspektivlər: Bu məsələ yerli kontekstdə xüsusilə əhəmiyyətlidir. Azərbaycan dili məhdud resurslara malik bir dildir; başqa sözlə, Azərbaycan dilindəki mətnlər böyük dil modellərinin (LLM) tədris materiallarında az təmsil olunur. Bu məsələ yerli fəlsəfi ədəbiyyatda da qeyd olunmuşdur. Quliyeva (2021) süni intellektin dil qabiliyyətinin miqyasını qiymətləndirərkən, Azərbaycan dili nümunəsində vurğulayır ki, "süni intellekt yalnız bu məqalənin müəllifinin ana dilində danışa bilməsi üçün ona 3 milyon söz-kəlmə lazımdır" [1, "Müzakirə" bölməsi]. Bu müşahidə az resurslu dillərin, o cümlədən Azərbaycan dilinin böyük dil modelləri qarşısındakı struktur dezavantajını konkret şəkildə göstərir və mənim paylanmış linqvistik idrak konsepsiyamın yerli kontekstdə niyə xüsusi aktualıq kəsb etdiyini əsaslandırır. Dilin bu cür məhdud təmsil olunması isə həm çətinlik, həm də fürsət yaradır. Problem ondadır ki, LLM-lər Azərbaycan dilində danışarkən daha çox səhv edirlər, dilin incəliklərini anlama bilmirlər və hətta mövcud olmayan sözləri "uydururlar". Lakin fürsət ondadır ki, bu vəziyyət yerli dilçilərə və filosoflara öz dilləri ilə bağlı yeni epistemoloji suallar qaldırmağa imkan verir.

Bu mövzu hələlik Azərbaycan akademik ictimaiyyəti daxilində kifayət qədər araşdırılmayıb. Əvvəlki tədqiqatımda [4, 5] süni intellektin sosial və ontoloji ölçülərinə diqqət yetirmişəm; bu məqalə həmin işin linqvistik-epistemoloji davamı kimi xidmət edir. Bu baxımdan, hesab edirəm ki, daha çox yerli tədqiqatlara, xüsusən də Azərbaycan dilinin dil öyrənmə mühitində təmsil olunması ilə bağlı empirik tədqiqatlara ehtiyac var. Bundan əlavə, inanıram ki, Şərq fəlsəfi ənənəsi (xüsusən də klassik İslam fəlsəfəsindəki "elm" anlayışı) müasir süni intellektin fəlsəfəsinə töhfə verə bilər; bu, mənim gələcək tədqiqat istiqamətlərimdən biridir.

Nəticə

Bu məqalədə böyük dil modellərində linqvistik bilik problemi bir-birini tamamlayan bir neçə nəzəri səviyyədə araşdırılmışdır. Tədqiqatın əsas məqsədi böyük dil modellərinin dil istehsalını yalnız texniki mexanizm kimi deyil, həm də epistemoloji və fəlsəfi problem kimi təhlil etmək olmuşdur. Bu çərçivədə dil və bilik münasibəti, anlama və simulyasiya mübahisəsi, insan–maşın qarşılıqlı əlaqəsində mənanın yaranması və bu prosesin müəlliflik, məsuliyyət və yerli dil konteksti baxımından doğurduğu nəticələr nəzərdən keçirilmişdir. Məqalənin ümumi məntiqi göstərir ki, böyük dil modelləri haqqında müzakirə yalnız “anlayır, yoxsa təqlid edir?” sualı ilə məhdudlaşdırıla bilməz.

Tədqiqatın başlanğıc nöqtəsi dil və bilik arasındakı klassik epistemoloji əlaqə olmuşdur. Vitgenşteynin mənanın istifadədə yaranması fikri ilə Çomskinin səriştə və performans fərqi böyük dil modellərinin mahiyyətini anlamaq üçün mühüm nəzəri zəmin yaradır. Çünki bu modellər məhz istifadə ilə səriştə arasındakı problemli sahədə yerləşir: onlar dili istifadə edir kimi görünür, lakin bu istifadənin arxasında insan şüuruna xas niyyət, inanc və doğruluq öhdəliyinin olub-olmaması mübahisəli qalır. Modellərin texniki əsasında növbəti tokenin ehtimalını hesablamaq mexanizmi dayansa da, miqyas böyüdükcə bu sadə mexanizmdən mürəkkəb və bəzən yaradıcı görünən linqvistik nəticələr meydana çıxır.

Bu baxımdan Searle, Bender, Floridi və funksionalist yanaşmalar arasında aparılan müzakirə göstərir ki, böyük dil modellərini yalnız bir tərəfdən izah etmək kifayət deyil. Searle-nin “sintaksis semantikaya gətirib çıxarmır” tezi modellərin insan mənasına və şüurlu niyyətə malik olmadığını vurğulamaq baxımından əhəmiyyətliyə. Benderin “stoxastik tutuquşu” metaforası da modellərin statistik təbiətini düzgün göstərir. Lakin bu yanaşmalar modellərin real istifadə prosesində yaratdığı interaktiv, kontekstual və integrativ imkanları tam izah etmir. Böyük dil modelləri sadəcə əvvəlki ifadələri təkrarlamır; onlar istifadəçi ilə qarşılıqlı əlaqədə yeni mətnlər və mənalər qurulmasında iştirak edir.

Məqalənin əsas nəticəsi ondan ibarətdir ki, böyük dil modelləri nə tam müstəqil anlama agentləri, nə də sadə mexaniki təkrar sistemləri kimi başa düşülməlidir. Daha uyğun yanaşma onları paylanmış linqvistik idrak müstəvisində izah etməkdir. Bu yanaşmaya görə, məna yalnız modelin daxilində və ya yalnız istifadəçinin şüurunda yaranmır. Məna istifadəçinin sualı, modelin cavabı, kontekstin qurulması, cavabın yoxlanılması, seçilməsi və yenidən formalaşdırılması prosesində meydana çıxır. Buna görə əsas sual modellərin insan kimi anlayıb-anlamaması deyil, onların insanla birlikdə hansı linqvistik və epistemik təcrübə yaratmasıdır.

Bu nəticə müəlliflik, bilik istehsalı və epistemik məsuliyyət anlayışlarının da yenidən nəzərdən keçirilməsini tələb edir. Böyük dil modelləri ilə yaradılan mətn çox vaxt tək bir subyektin məhsulu olmur: istifadəçi sualı formalaşdırır, model cavab yaradır, istifadəçi həmin cavabı dəyişdirir, qəbul edir və ya rədd edir. Nəticədə ortaya çıxan mətn nə tam şəkildə istifadəçiyə, nə də tam şəkildə modelə aid edilə bilər. Bu isə fərdi müəlliflik modelindən paylanmış müəlliflik modelinə keçidi aktuallaşdırır. Lakin müəllifliyin paylanması məsuliyyətin itməsi anlamına gəlməməlidir. Əksinə, belə şəraitdə epistemik və normativ məsuliyyət daha aydın müəyyən edilməlidir.

Azərbaycan dili kontekstində bu məsələ xüsusi əhəmiyyət daşıyır. Böyük dil modellərinin yaratdığı problemlər yalnız qlobal texnoloji mühitlə bağlı deyil, həm də yerli dil, akademik yazı ənənəsi və institusional praktika ilə əlaqəlidir. Azərbaycan dili az resurslu dil kimi modellərdə daha zəif təmsil oluna bilər və bu, kontekstin düzgün tutulmaması, semantik incəliklərin itməsi və səhv cavabların artması kimi risklər yaradır. Buna görə gələcək tədqiqatlarda paylanmış linqvistik idrak konsepsiyasının empirik əsaslandırılması, real insan–model dialoqlarının təhlili və Azərbaycan dilində böyük dil modellərinin fəaliyyət xüsusiyyətlərinin ayrıca araşdırılması vacib istiqamətlər kimi qalır. Yekun olaraq, böyük dil

modellərini nə düşünən insan subyekti kimi şışirtmək, nə də sadəcə statistik aldanış kimi rədd etmək doğrudur; onları insanla qarşılıqlı təsirdə mənə yaradan hibrid linqvistik mühit kimi başa düşmək daha məhsuldar fəlsəfi yanaşmadır.

ƏDƏBİYYAT

1. Quliyeva, X. (2021, 10 iyun). Süni intellekt fəlsəfəsinin progressiv və regressiv inkişaf məsələləri. AMEA. <https://science.gov.az/az/news/open/17146>
2. Mahmudov, M. (2024). Süni intellektin linqvistik problemləri. Bakı: Elm və təhsil.
3. Məmmədov, Ə., & Qurbanov, F. (2019). Lütfi Zadənin Qeyri-səlis çoxluqlar nəzəriyyəsinin məntiqi-qnoseoloji təhlili. *Metafizika*, 2(1), 7–29.
4. Mustafayev, A. (2025a). İnsan və süni intellekt: Ontoloji münaqişələr. *Şərq Fəlsəfəsi Problemləri Jurnalı*, (33), 25–36. <https://doi.org/10.59849/2219-7370.2025.33.25>
5. Mustafayev, A. (2025b). Birinci Dünya Müharibəsindən müasir bilik cəmiyyətinə: Texnoloji və informasiya determinizminin təkamülü. *Akademik Tarix və Düşüncə Dergisi*, 12(5), 296–306. <https://dergipark.org.tr/tr/pub/atdd/issue/94724/1800680>
6. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
7. Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On meaning, form, and understanding in the age of data. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 5185–5198. <https://doi.org/10.18653/v1/2020.acl-main.463>
8. Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., et al. (2021). On the opportunities and risks of foundation models. *Stanford Center for Research on Foundation Models*. <https://arxiv.org/abs/2108.07258>
9. Chomsky, N. (1957). *Syntactic Structures*. The Hague: Mouton.
10. Chomsky, N. (1965). *Aspects of the Theory of Syntax*. Cambridge, MA: MIT Press.
11. Clark, A., & Chalmers, D. (1998). The extended mind. *Analysis*, 58(1), 7–19. <https://doi.org/10.1093/analys/58.1.7>
12. Coeckelbergh, M. (2022). Democracy, epistemic agency, and AI: Political epistemology in times of artificial intelligence. *AI and Ethics*, 3(4), 1341–1350. <https://doi.org/10.1007/s43681-022-00239-4>
13. Dennett, D. C. (2017). *From Bacteria to Bach and Back: The Evolution of Minds*. New York: W. W. Norton & Company.
14. Floridi, L. (2023). AI as agency without intelligence: On ChatGPT, large language models, and other generative models. *Philosophy & Technology*, 36(1), 1–7. <https://doi.org/10.1007/s13347-023-00621-y>
15. Hutchins, E. (1995). *Cognition in the Wild*. Cambridge, MA: MIT Press.
16. Mitchell, M., & Krakauer, D. C. (2023). The debate over understanding in AI's large language models. *Proceedings of the National Academy of Sciences*, 120(13), 1–13. <https://doi.org/10.1073/pnas.2215907120>
17. Popper, K. R. (1972). *Objective Knowledge: An Evolutionary Approach*. Oxford: Clarendon Press.
18. Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–457. <https://doi.org/10.1017/S0140525X00005756>
19. Shanahan, M. (2024). Talking about large language models. *Communications of the ACM*, 67(2), 68–79. <https://doi.org/10.1145/3624724>

20. Wittgenstein, L. (1922). Tractatus Logico-Philosophicus (C. K. Ogden, Trans.). London: Kegan Paul, Trench, Trubner & Co.
21. Wittgenstein, L. (1953). Philosophical Investigations (G. E. M. Anscombe, Trans.). Oxford: Basil Blackwell.

УДК 1:001:004.8:81

*Айхан Адалет оглы Мустафаев
Аспирант Института Философии и Социологии
Национальной Академии Наук Азербайджана
Физули Мухаммед оглы Гурбанов,
Мухаммед Султан оглы Джабраилов*

БОЛЬШИЕ ЯЗЫКОВЫЕ МОДЕЛИ И ПРОБЛЕМА ЛИНГВИСТИЧЕСКОГО
ЗНАНИЯ: ПОНИМАНИЕ ИЛИ СИМУЛЯЦИЯ
РЕЗЮМЕ

Ключевые слова: большие языковые модели, Лингвистические знания, искусственный интеллект, моделирование, эпистемология, распределенное лингвистическое познание, общество знаний

В статье рассматривается проблема лингвистического знания в больших языковых моделях с эпистемологической и философской точек зрения. Основная цель исследования заключается в анализе вопроса о том, следует ли рассматривать способность таких систем, как ChatGPT, Claude, Gemini и других больших языковых моделей, создавать тексты на человеческом языке как подлинное понимание или как простую статистическую симуляцию. В статье утверждается, что способность таких систем генерировать связные и внешне осмысленные тексты не доказывает наличия у них понимания языка в человеческом смысле этого слова. Вместе с тем сведение больших языковых моделей исключительно к “стохастическим попугаям” также не раскрывает всей сложности данного эпистемологического феномена. Опираясь на идеи Витгенштейна, Серла, Хомского, Флориди, Бендер и других исследователей, автор предлагает третью позицию, выраженную в концепции “распределённого лингвистического познания”. Согласно данной концепции, значение не находится исключительно ни в человеческом сознании, ни внутри алгоритмической модели, а формируется в процессе взаимодействия человека и машины. В статье также анализируются проблемы галлюцинаций, эпистемическая ответственность, авторство, а также риски и возможности, возникающие в контексте малоресурсных языков, включая азербайджанский язык. В заключении подчеркивается, что большие языковые модели требуют переосмысления статуса языка, значения и знания в обществе знаний.

*Ayxan Adalet Mustafayev
PhD Candidate, Institute of Philosophy and Sociology,
Azerbaijan National Academy of Sciences
Fuzuli Mahammad Gurbanov,
Mahammad Sultan Jabrayilov*

LARGE LANGUAGE MODELS AND THE PROBLEM OF LINGUISTIC
KNOWLEDGE: UNDERSTANDING OR SIMULATION
SUMMARY

Keywords: Large language models, linguistic knowledge, artificial intelligence, simulation, epistemology, distributed linguistic cognition, knowledge society

The article examines the problem of linguistic knowledge in large language models from an epistemological and philosophical perspective. Its main aim is to analyze whether the ability of systems such as ChatGPT, Claude, Gemini, and other LLMs to generate human-like linguistic output should be understood as genuine understanding or merely statistical simulation. The article argues that the production of coherent and seemingly meaningful texts does not prove that these systems understand language in the human sense. At the same time, reducing LLMs to mere “stochastic parrots” oversimplifies the epistemological significance of this phenomenon. Drawing on the theoretical perspectives of Wittgenstein, Searle, Chomsky, Floridi, Bender, and other scholars, the author proposes a third position through the concept of “distributed linguistic cognition.” According to this approach, meaning is located neither solely in the human mind nor entirely within the algorithmic model; rather, it emerges through human–machine interaction. The article also discusses the problems of hallucinations, epistemic responsibility, authorship, and the risks and opportunities created by LLMs in the context of low-resource languages such as Azerbaijani. The conclusion emphasizes that large language models require a reconsideration of the status of language, meaning, and knowledge in the knowledge society.

Daxil oldu: 22.06.2026-cı il