

DOI: <https://doi.org/10.36719/2789-6919/53/137-142>

Aynurə Ələkbərova
Bakı Mühəndislik Universiteti
<https://orcid.org/0009-0002-5443-1255>
aelekberova@beu.edu.az

Süni intellektin potensialı və təhlükələri

Xülasə

Məqalədə ənənəvi süni intellektdən (Sİ) generativ süni intellektə keçid mərhələləri qısa şəkildə şərh olunur. Generativ Sİ-in insanlara verdiyi bir çox faydalarla yanaşı, onun yaratdığı təhlükələrdən də bəhs olunur. Eləcə də göstərilir ki, Sİ-in proqnozlaşdırma imkanları yaxşılaşdıqca, sərvətin və təsirin bölüşdürülməsi mühakimə keyfiyyətindən getdikcə daha çox asılı olacaqdır. Yüksək və aşağı ixtisaslı işçilər arasındakı fərqi, işin proqnozlaşdırıcı aspekti ilə müəyyən edildiyi sahələrdə Sİ daha aşağı ixtisaslı işçilərə fayda verəcək, çünki Sİ proqnozlaşdırmaqda insanları əvəz edə bilər. Bu, müəyyən bir sənayedə işçilər arasında əmək məhsuldarlığında fərqləri və buna görə də gəlir bərabərsizliyini hamarlaşdıracaq və zaman keçdikcə aşağı əmək haqqı olan ölkələrdə, hətta daha aşağı bacarıqlara malik olanlar üçün də əmək haqqının artmasına gətirib çıxaracaqdır.

Açar sözlər: *süni intellekt, kompüter, program, şəbəkə, proqnozlaşdırma*

Aynura Alekberova
Baku Engineering University
<https://orcid.org/0009-0002-5443-1255>
aelekberova@beu.edu.az

The Potential and Dangers of Artificial Intelligence

Abstract

The article briefly describes the stages of the transition from traditional artificial intelligence (AI) to generative AI. It discusses the many benefits that generative AI brings to humans, as well as the dangers it poses. It also shows that as AI's predictive capabilities improve, the distribution of wealth and influence will increasingly depend on the quality of judgment. In areas where the difference between high- and low-skilled workers is determined by the predictive aspect of the job, AI will benefit lower-skilled workers because AI can replace humans in making predictions. This will smooth out differences in labor productivity and, therefore, income inequality between workers in a given industry, and over time, will lead to higher wages in low-wage countries, even for those with lower skills.

Keywords: *artificial intelligence, computer, program, network, forecasting*

Giriş

Təsəvvür edin ki, kompüter sistemləri rəssamlar kimi, alimlər kimi, iqtisadçılar kimi insan intellektini təqlid edən informasiya məhsulları yaradırlar. Kompüter elmləri üzrə dünya şöhrətli Alan Turing kompüterlərin bu bacarıq səviyyəsindən ilk dəfə 1950-ci ildə öz məqalələrində söhbət açmışdı. Chat, GPT və digər sözdə “generativ süni intellekt alətləri” sayəsində onun proqnozlaşdırdığı “Təqlid oyunu” indi reallığa çevrilib. Sanki bir zamanlar yalnız elmi fantastikada mövcud olan bir kainata daxil olmuşuq. Bəs generativ süni intellekt (Sİ) nədir?

Generativ Sİ bu günə qədər maşın öyrənməsində ən maraqlı irəliləyişdir. Bu, Sİ-in mürəkkəb məlumat strukturlarını insan kimi anlamaq və onlarla qarşılıqlı əlaqə qurmaq qabiliyyətində böyük bir sıçrayışı təmsil edir və yaradıcılıq və məhsuldarlıqda artım yaratmaq potensialına malikdir.

Lakin bu, bəşəriyyət üçün vacib suallar doğurur. İnnovasiyanın əsas mərhələləri onun hazırkı inkişaf vəziyyətinə gedən yolu artıq formalaşdırmışdır.

Tədqiqat

1960-cı illərdə yaradılan ELIZA adlı proqram, ona verilən sualları insan kimi cavablandırmaq qabiliyyəti ilə alimləri heyran etdi. Proqram çox mürəkkəb deyildi və müəyyən qaydalara əsaslanırdı. Həmin proqram bu gün bildiyimiz chatbotların xəbərçisi idi. İki onillikdən sonra süni neyron şəbəkələri ortaya çıxdı. İnsan beynindən fəaliyyətinə əsaslanaraq modelləşdirilən bu şəbəkələr, kompüterlərə dilin nüanslarını anlamaq və şəkilləri tanımaq kimi yeni bacarıqlar verdi. Bununla belə, məhdud təlim məlumatları və qeyri-kafi hesablama gücü real irəliləyişə mane oldu. Maraqlıdır ki, bu iki resurs hər il iki dəfə artmağa davam edərək 2000-ci illərdə süni intellektin üçüncü dalğası üçün zəmin yaratdı: dərin öyrənmə (Ajay, Gans, Goldfarb, 2018).

Google Translate kimi yeniliklər, Alexa və Siri kimi rəqəmsal köməkçilər və özünü idarə edən avtomobillərin meydana gəlməsi sayəsində kompüterlər dünyanı anlamağa, eləcə də digər sistemlərlə və insanlarla qarşılıqlı əlaqə yaratmağa başladı. Bütün bu irəliləyişlərə baxmayaraq, tam mənzərə hələ də yerində deyildi. Maşınlar kömək edə və proqnozlaşdırırdı, lakin onlar insan ünsiyyətinin incəliklərini həqiqətən dərk edə bilmirdilər və insanlar kimi informasiya məhsulları yaratmaqda zəif idilər.

Daha sonra, 2014-cü ildə generativ rəqib şəbəkələr (GAN) bir-birinin bacarıqlarını davamlı olaraq təkmilləşdirmək üçün rəqabət aparan iki neyron şəbəkənin bacarıqlarından istifadə etdi. "Generator" simulyasiya edilmiş məlumat, mətn və ya şəkillər yaratdı, real informasiya məhsulunu simulyasiya edilmiş məhsuldan ayırmağa çalışdı. İki şəbəkə arasındakı bu rəqabət süni intellektin mürəkkəb strukturları başa düşməsi və təkrar istehsalında inqilab etdi. Getdikcə artan məlumat və hesablama gücü ilə dəstəklənən GAN-lar və diqqət mexanizmləri indiyə qədər ən heyvətəmiz chatbot olan ChatGPT üçün əsas yaratdı. O, "OpenAI" tərəfindən 2022-ci ilin noyabrında istifadəyə verildi və tezliklə digər böyük texnoloji firmalar generativ süni intellektə əsaslanan öz chatbotlarını təqdim etdilər (Brynjolfsson, Li, Raymond, 2023).

Təbii ki, Sİ iqtisadiyyat və maliyyədə yeni bir anlayış deyil. Ənənəvi süni intellekt formaları (qabaqcıl analitika, maşın öyrənməsi, proqnozlaşdırıcı dərin öyrənmə) hələlik bazar tendensiyyalarını qiymətləndirmək və maliyyə məhsullarını uyğunlaşdırmaq üçün istifadə olunur. Generativ Sİ-i fərqləndirən əsas cəhətlərdən biri - onun mürəkkəb məlumatları daha yaradıcı şəkildə təhlil və şərh etmək qabiliyyətidir. İqtisadi göstəricilər və ya maliyyə dəyişənləri arasındakı mürəkkəb əlaqələri təhlil edərək, o, təkcə proqnozlar deyil, həm də sənayenin fəaliyyət tərzini əhəmiyyətli dərəcədə dəyişdirə biləcək alternativ ssenarilər, illüstrativ qrafiklər və hətta kod parçaları (yəni kompüter proqramları) da hazırlayır (Ball, 2023).

Ənənəvi süni intellektdən generativ süni intellektə keçid həm dövlət, həm də özəl sektorda yeni imkanlar açdı. Hökumətlər vətəndaşlara daha yaxşı xidmət göstərmək və işçi qüvvəsi çatışmazlığını aradan qaldırmaq üçün bu daha ağıllı vasitələrdən istifadə etməyə başlayır. Mərkəzi bankların iqtisadi proqnozları dəqiqləşdirmək və fırladaçılıq da daxil olmaqla riskləri daha yaxşı izləmək və minimuma endirmək üçün generativ süni intellektin böyük həcmdə bank məlumatlarını təhlil etmək potensialı olmalıdır.

İnvestisiya şirkətləri səhm qiymətlərində və dünya bazarı səviyyəsində incə dəyişiklikləri aşkar etmək üçün generativ süni intellektə müraciət edir, daha yaradıcı variantlar təklif etmək üçün daha böyük bilik hovuzundan (verilənlər bazasından (VB)) istifadə edir, potensial olaraq daha sərfəli investisiya strategiyalarını həyata keçirmək üçün imkanlar yaradır. Eyni zamanda, sığorta şirkətləri fərdi ehtiyac və seçimlərə daha yaxından uyğun gələn fərdiləşdirilmiş siyasətlər yaratmaq məqsədi ilə generativ modellərin potensialını araşdırırlar (Begley, Ioannidis, 2015).

Generativ süni intellekt inanılmaz sürətlə irəliləyir, süni intellektin iqtisadiyyat və maliyyə sahəsində edə biləcəklərinin sərhədlərini aşır və həll olunmuş köhnə problemlərə yeni həllər təklif edir. Bəzi insanlar bu inqilabi sıçrayışa şübhə ilə yanaşırlar. Onlar düşünürlər ki, süni intellekt stoxastik tutuquşu kimi mənasız və yalan faktlar, "halüsinasiyalar" adlı fenomen yarada bilər. Onlar bu sözlərin nə demək olduğunu həqiqətən başa düşmür və qeyd edirlər ki, Chat GPT-nin bilikləri ən

son təlim tarixi ilə məhdudlaşır. Bəs bu argumentlər innovasiyanın yüksək inkişaf sürətinə qarşı nə qədər davam edəcək?

Bununla belə, generativ süni intellekt ətrafında yaranan ilk həyəcan öz yerini həqiqi şəkildə artan narahatlıqlara verdi. Təlim məlumatlarında mövcud qərəzlərin gücləndirilməsi və ya qərarlarda şəffaflığın olmaması kimi Sİ ilə bağlı ənənəvi narahatlıqlar yenidən aktualıq qazanmağa başlamışdır. Bu narahatlıqlar - Sİ-in silahlanmada tətbiqidir.

Xüsusilə narahatlıq doğuran risklərdən biri generativ süni intellektin insanların əvvəlcədən mövcud olan inancları və baxışları ilə sıx bağlı olan, əks-səda kameralarını potensial olaraq gücləndirən hekayələr söyləmək qabiliyyətidir. Pis aktyorlar bu qabiliyyətdən yazılı sözlərdən kənarda istifadə edə bilirlər: 2022-ci ilin martında Ukrayna prezidenti Volodimir Zelenskinin Rusiya qüvvələrinə təslim olmasını göstərən süni intellekt əsasında hazırlanmış video yayımlandı. Bu kimi hadisələr generativ süni intellektin siyasəti, bazarları və ictimai rəyi manipulyasiya etmək üçün necə silahlandırıla biləcəyini nümayiş etdirir.

İstər uydurulmuş hekayə, istərsə də düzəldilmiş şəkil və ya sintez edilmiş video olsun, generativ Sİ yaradıcılığı o qədər inandırıcı ola bilər ki, yalan - realıq hissi yaradar. Bu dezinformasiya yayıla bilər, çaxnaşmaya səbəb ola bilər və hətta iqtisadi və ya maliyyə sistemlərini görünməmiş səmərəlilik və intensivliklə destabilizasiya edə bilər. Bu, həmişə qəsdən olmaya bilər: maşınlar halüsinasiyalar vasitəsilə istəmədən dezinformasiya yayırlar. Bu da insanlar və ya dövlətlər arasında münaqişələrə səbəb ola bilər, fəlakətlər törədə bilər və s.

Süni intellektin yaratdığı təhlükə yalnız manipulyasiya ilə məhdudlaşmır. İş yerlərinin məhv edilməsi başqa bir narahatlıq doğurur, çünki generativ süni intellekt irəliləməyə davam edir, insanlar tərəfindən əvvəllər yerinə yetirilən vəzifələri potensial olaraq avtomatlar tutur. Bu da bir çox iş yerlərinin itirilməsinə səbəb olur və yenidən bacarıq və yenidən təlim strategiyalarının yaradılmasını tələb edir.

Bu ilin əvvəlində aparıcı süni intellekt mütəxəssisləri, o cümlədən ChatGPT-nin yaradıcısı, “Pandemiya və nüvə müharibəsi kimi digər sosial risklərlə yanaşı, süni intellektin məhv edilməsi riskinin azaldılması qlobal prioritet olmalıdır” deyə xəbərdarlıq məktubu imzaladı. Onlar Turinqin onilliklər əvvəl qaldırdığı narahatlıqları təkrarladılar, o, “maşınların nəhayət həyatımızı nəzarətə götürməsi təhlükəsi var” deyə xəbərdarlıq etdilər.

Biz texnologiya və etikanın kəsişməsində dayanırıq. Generativ Sİ, böyük vədləri və dərin ekzistensial sualları ilə geri qaytarıla bilməz. Onun dəyişdirici gücündən istifadə edərkən, Turinqin davamlı məsləhətlərini xatırlamaq çox vacibdir. Generativ süni intellekt ayıq nəzarət, yeni tənzimləyici çərçivələr və ümumbəşəri dəyərlərə uyğun gələn etik, şəffaf, hesabatlı innovasiyalara sarsılmaz sadıqlıq tələb edən monumental dəyişiklikdir (Chubb, Cowling, Reed, 2022).

Milyonlarla dahinin yaşadığı bir ada təsəvvür edin. Onlar kompüterin edə biləcəyi hər şeyin mütəxəssisidir. Onlar gecə-gündüz işləyirlər və az maaşla işləməkdən məmnundurlar. İndi təsəvvür edin ki, bu insanlar qlobal iqtisadiyyata daxil olanda nə qədər böyük suallar yaranacaqdır.

Onların iştirakı bazarlara, maaşlara və güc münasibətlərinə necə təsir edəcək? Bu dahilərin meydana çıxması bolluq və firavanlığa və ya həddindən artıq qeyri-sabitliyə gətirib çıxara bilər və bu, ən azından bizim bu münasibətlərlə necə davranacağımızdan asılı olacaqdır.

Yeni bolluq dövründə məhsuldarlıq və iqtisadi artım yüksələcək, sosial rifah çiçəklənəcəkdir. Bu yeni işçi qüvvəsi bənzərsiz zəkaları ilə səhiyyədən təhsilə, texnologiyaya qədər sənayedə inqilab edəcəkdir. Ofis tapşırıqları tez və dəqiq yerinə yetiriləcək, daha mənalı fəaliyyətlər üçün vaxt ayrılacaqdır. Bir çox xidmətlərin qiyməti aşağı düşəcək, həyat səviyyəsi yüksələcəkdir.

Qeyri-sabit ssenari nə kimi görünür? Dahilər eyni işi müqayisə olunmayacaq qədər az maaşla görsələr, bilik işçiləri və mütəxəssislər kütləvi işsizliklə üzləşəcəklər. Aşağı əmək haqqı və iş yerlərini itirmək qorxusu müxtəlif sənaye sahələrinə zərbə vuracaq, bu da əhalinin orta sinfinin kəskin azalmasına və bərabərsizliyin kəskinləşməsinə gətirib çıxaracaqdır. Bu dahilər bir neçə şirkət və ya dövlət tərəfindən ələ keçirilərsə, sərvət və nüfuzun görünməmiş inhisarlaşması başlaya bilər ki, bu da daha kiçik şirkətləri və iqtisadiyyatı zəif olan ölkələri tənəzzül sərhədlərinə itələyər. Bu, innovasiya prosesini ləngidə və qlobal ziddiyyətləri kəskinləşdirə bilər.

İntellektual və praktik sferalarda dahilər hökmranlıq edərsə, insanın fərdiliyi və yaradıcılığı dəyərini itirə bilər. Çoxlarının özünü lazımsız hiss etdiyi bir şəraitdə cəmiyyətlər “biz kimik?” ekzistensial sualları üzərində əzab çəkəcəklər və geniş iğtişaslara səbəb ola biləcəklər. Beləliklə, bu dahilər iqtisadi xaosa, sosial qarışıqlığa və dünyanın bərabərsizliyə qərq olmasına səbəb ola bilər (Hanson, Stall, Cutcher-Gershenfeld, Vrouwenvelde, Wirz, Rao, Peng, 2023).

Bu dahilər adası haqqında düşünməyə dəyər, çünki getdikcə daha çox mütəxəssislər hesab edirlər ki, biz məhz belə bir texnoloji sıçrayış astanasında ola bilərik. Məsələn, 2023-cü ildə süni intellekt şəbəkələri sahəsindəki kəşflərinə görə 2024-cü ildə Nobel mükafatı alan və dünyada “Süni intellektin atası” ünvanını alan kanadalı ingilis Ceffri Hinton (Geoffrey Hinton) demişdir ki, süni intellekt yaxın 5-20 il ərzində insan intellektini keçə bilər. Bəzi ekspertlər bunun daha tez baş verə biləcəyini iddia edirlər.

İnsan intellektini üstələyən süni intellektin yaranmasının rifahın artmasına və ya qeyri-sabitliyə gətirib çıxarıb-çıxarmayacağı, çox güman ki, onun bərabərsizliyə necə təsir etməsindən asılı olacaqdır. 1960-cı illərin kompüter inqilabından bəri, Nobel mükafatı laureatı, Massaçusset Texnologiya İnstitutunun professoru, 1967-ci ildə İstanbulda doğulmuş Daron Acemoğlu da daxil olmaqla bir çox iqtisadçı iddia edirdi ki, texnoloji tərəqqi yüksək ixtisaslı və təcrübəli işçilərə tələbi artırmaqla və aşağı ixtisaslı işçi qüvvəsinə tələbi azaltmaqla gəlir bərabərsizliyini pisləşdirə bilər ki, bu da “bacarıq meyli” kimi tanınan bir fenomendir. Son iki araşdırma, bacarıq meyli konsepsiyasının Sİ inqilabına necə tətbiq oluna biləcəyinə işıq salır (Toni, 2024).

Stanford Universitetinin professoru, Sİ-in biznes strategiyasına təsiri sahəsində elmi işlər aparan Erik Brinolfson (Erik Brynjolfsson) və həmkarlarının zəng mərkəzi işçiləri haqqında məlumatları təhlil etdiyi başqa bir araşdırma, daha az ixtisaslı işçilərin süni intellektdən qeyri-mütənasib şəkildə faydalandığını göstərir. Təcrübəli və yüksək ixtisaslı agentlərin məhsuldarlığı faktiki olaraq dəyişməz qalsa da, yeni və daha az ixtisaslı işçilərin məhsuldarlığı 34 faiz artır. Müəlliflər, xüsusilə süni intellekt alətlərinin məhsuldarlığı (saatda cavablandırılan sualların sayı ilə ölçülür) orta hesabla 14 faiz artırdığını müəyyən etmişlər. Məsələn, Sİ, daha bacarıqlı həmkarlarının müəyyən bir problemi necə həll edəcəyini proqnozlaşdırmaqla daha az bacarıqlı agentlərin işini yaxşılaşdırma bilər. Bu şəraitdə Sİ gəlir bərabərsizliyini azaldır (Mauro, Jaumotte, Li, Melina, Panton, Pizzinelli, Rockall, Tavares, 2024).

Niyə bəzi tədqiqatçılar süni intellektin daha az ixtisaslı işçilərə digərlərindən daha çox kömək etdiyini tapırlar? Tələbə debatçıları ilə zəng mərkəzi operatorları arasında fərq nədir? Biz bunun mühakimə (qərar vermədə əsas komponentdir) və proqnozlaşdırma ilə əlaqəli olduğuna inanırıq. Bu iki aspekt müxtəlif nəticələrə müəyyən ehtimal (proqnozlaşdırma) və onların nəticələrinə müəyyən qiymət (mühakimə) təyin edən tətbiqi ehtimal nəzəriyyəsinin bir qolu olan qərarlar nəzəriyyəsində əsas rol oynayır (Fabrizio and oth., 2023).

1983-cü ildə anadan olmuş ispan iqtisadçısı Roldan Monesin tədqiqatının nəticələri göstərir ki, yüksək səriştəli və aşağı səriştəli debatçılar arasındakı bərabərsizlik mühakimə ilə bağlı ola bilər. Müəllif belə nəticəyə gəlir ki, süni intellektdən istifadə edən yüksək səriştəli debatçılar inandırıcılıq, natiqlik və əks arqumentlər üzrə daha yüksək nəticələr alırlar. Lakin onların təqdimatın aydınlığına görə balları artmayıb. Bu, hazır cavab əvəzinə Sİ alətinin seçimlər təklif etdiyini və daha yüksək səviyyəli səriştəli iştirakçıların ən perspektivli arqumentləri seçmək vəzifəsinin öhdəsindən daha yaxşı gəldiyini göstərir. Nəticədə, əsas üstünlükləri daha yüksək səviyyəli səriştəli iştirakçılar əldə etdilər ki, bu da bacarıq bərabərsizliyini artırdı (Aidan, 2024).

Brinolfsonun araşdırmasında, az və ya çox ixtisaslı zəng mərkəzi operatorları arasındakı əsas fərq, müştərinin ən yaxşı şəkildə qane edəcək cavabı proqnozlaşdırmaq bacarığı idi. Süni intellekt isə bu vəzifənin öhdəsindən bacarıqlı işçilərdən heç də pis gəlmirdi. Müxtəlif növ xətalara bağlı nisbi məsrəflərin təhlilinə gəlincə, onun əhəmiyyəti o qədər də yüksək deyildi, çünki belə nəticələri düzgün qiymətləndirmək imkanı azdır (Bashirov, 2002).

Süni intellektin proqnozlaşdırma imkanları yaxşılaşdıqca, sərvətin və təsirin bölüşdürülməsi mühakimə keyfiyyətindən getdikcə daha çox asılı olacaqdır. Yüksək və aşağı ixtisaslı işçilər arasındakı fərqlin, işin proqnozlaşdırıcı aspekti ilə müəyyən edildiyi sahələrdə Sİ daha aşağı ixtisaslı

işçilərə fayda verəcək, çünki Sİ proqnozlaşdırmaqda insanları əvəz edə bilər. Bu, müəyyən bir sənayedə işçilər arasında əmək məhsuldarlığında fərqləri və buna görə də gəlir bərabərsizliyini hamarlaşdıracaq və zaman keçdikcə aşağı əmək haqqı olan ölkələrdə, hətta daha aşağı bacarıqlara malik olanlar üçün də əmək haqqının artmasına gətirib çıxaracaq (Vasi'ev, Tyrov, Vladzimirskii, Shul'kin, Arzamasov, 2024). Məsələn, Hindistanda zəng mərkəzi və inzibati maaşlar ABŞ-a nisbətən arta bilər... (Dünyanın ən iri şirkətlərinin əsas proqramçıları hindistanlılardır !).

Lakin yüksək və aşağı bacarıqlar arasındakı fərqin mühakimə keyfiyyəti ilə müəyyən edildiyi sahələrdə Sİ, ilk növbədə daha yüksək ixtisaslı işçilərə fayda verəcəkdir (Vasilev, Tyrov, Vladzimirskyy, Shulkin, Arzamasov, 2024). Nəticə etibarilə, bu sənayelərdə məhsuldarlıq və gəlir fərqi daha da genişlənəcəkdir. İşçi qüvvəsi əvvəllər daha az cəlbədicə olan yüksək maaşlı bölgələrə köçməyə başlaya bilər, çünki, yüksək bacarıqların faydaları xərcləri doğrultmur. İstedadlı tələbələr getdikcə daha çox ABŞ universitetlərini seçdikcə və ABŞ-da yaşayan alimlər elmi kəşflərin, mükafatların, nəşrlərin və patentlərin sayına görə öz həmyaşıdlarını qabaqladıqları üçün, ABŞ-a innovasiya fəaliyyətinin axını sürətləndirilə bilər (Vasilev, Tyrov, Vladzimirskyy, Shulkin, Arzamasov, 2024; Zinchenko, Arzamasov, Kremneva, Vladzimirskii, Vasil'ev, 2023). Bu yaxınlarda “Süni intellektin atası” kimi tanınan 75 yaşlı Cefri Hinton süni intellekt sahəsindəki fəaliyyətlərinə və gördüyü işlərə görə peşman olduğunu bildirərək “Google”dan istefa etdi.

“New York Times” nəşrinə açıqlamalar edən Hinton, “Bu sahədəki inkişaflardan qaynaqlanan potensial təhlükələr üzündən süni intellektlə bağlı fəaliyyətimdən peşmanam” - deyərək, fikrini dilə gətirib. Süni intellektə əsaslanan çatbotların hazırda insanlardan daha təhlükəli olmadığını, lakin bəzi potensial təhlükələrinin olduqca qorxulu olduğunu söyləyən Hinton onların yaxın gələcəkdə insanlardan daha ağıllı olacağı qənaətinə gəldiyini ifadə edib.

Nəticə

Süni intellekt sürətlə inkişaf edir və yeni mərhələlərə qədəm qoyur. Lakin idarəetmə təcrübələri, infrastruktur, təhsil, tənzimləmə və istehlakçı tələbi kimi sahələr və parametrlər yavaş-yavaş dəyişir, ona görə Sİ-in inkişafından geri qalır. Lakin zaman keçdikcə bu geri qalma aradan qalxaraq global iqtisadiyyata əhəmiyyətli təsir göstərəcək və iqtisadi sabitlik bu keçidi necə idarə etməyimizdən asılı olacaqdır. Bu asılılığı zaman göstərəcəkdir. Beləliklə, bəşəriyyəti onlara veriləcək möhtəşəm faydalarla bərabər, potensial təhlükələr də gözləyir.

Ədəbiyyat

1. Ajay, A., Gans, J., Goldfarb, A. (2018). *Prediction Machines: The Simple Economics of Artificial Intelligence*. Harvard Business Review Press.
2. Erik, B., Li, D., Raymond, L. (2023). *Generative AI at Work*. National Bureau of Economic Research.
3. Ball, P. (2023). Is AI leading to a reproducibility crisis in science? *Nature*, 624(7990), 22-25.
4. Begley, C.G., Ioannidis, J.P. (2015). Reproducibility in science: improving the standard for basic and preclinical research. *Circulation research*, 116(1), 116-126.
5. Chubb, J., Cowling, P., Reed, D. (2022). Speeding up to keep up: exploring the use of AI in the research process. *AI & society*, 37(4), 1439-1457.
6. Cows, J., Tsamados, A., Taddeo, M., Floridi, L. (2023). The AI gambit: leveraging artificial intelligence to combat climate change – opportunities, challenges, and recommendations. *AI & Society*, 1-25.
7. Hanson, B., Stall, S., Cutcher-Gershenfeld, J., Vrouwenvelder, K., Wirz, C., Rao, Y., Peng, G. (2023). Garbage in, garbage out: mitigating risks and maximizing benefits of AI in research. *Nature*, 623(7985), 28-31.
8. Toni, R. (2024). *When GenAI Increases Inequality: Evidence from a University Debating Competition*. EsadeEcPol Working Paper, Esade Center for Economic Policy.

9. Mauro, G., Jaumotte, F., Li, L., Melina, G., Panton, A., Pizzinelli, C., Rockall, E., Tavares, M. (2024). *Gen-AI: Artificial Intelligence and the Future of Work*. IMF Staff Discussion Note 2024/001. International Monetary Fund.
10. Fabrizio, D. and others. (2023). *Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality*. Harvard Business School Working.
11. Aidan, T. (2024). *Artificial Intelligence, Scientific Discovery, and Product Innovation*. ArXiv preprint, Cornell University.
12. Bashirov, N. (2002). *Ehkspertnye sistemy*. Elm.
13. Vasi'ev, YU.A., Tyrov, I.A., Vladzimirskii, A.V., Shul'kin, I.M., Arzamasov, K.M. (2024). Avtonomnyi iskusstvennyi intellekt dlya sortirovki rezul'tatov profilakticheskikh luchevykh issledovaniy. *Profilakticheskaya meditsina*, 27(7): 23-29. <https://doi.org/10.17116/profmed20242707123>
14. Vasilev, Yu.A, Tyrov, I.A., Vladzimirskyy, A.V., Shulkin, I.M., Arzamasov K.M. (2024). Autonomous artificial intelligence for sorting the preventive imaging studies' results. *Russian Journal of Preventive Medicine*, 27(7):23-29. <https://doi.org/10.17116/profmed20242707123>
15. Zinchenko, V.V., Arzamasov, K.M., Kremneva, E.I., Vladzimirskii, A.V., Vasil'ev, YU.A. (2023). Tekhnologicheskie defekty programmnoy obespecheniya s iskusstvennym intellektom. *Digital Diagnostics*, 4(4), 593-604. DOI: <https://doi.org/10.17816/DD501759>

Daxil oldu: 07.09.2025

Qəbul edildi: 12.12.2025