



A Nemzeti Közsolgálati Egyetem Információs Társadalom Kutatóintézete havi hírlevele a mesterséges intelligencia alkalmazásáról, társadalmi hatásairól és kérdéseiről

2025 június

Az NKE ITKI honlapja: itki.uni-nke.hu

A hírlevél tartalma a Creative Commons Nevezd meg! – Ne add el! – Így add tovább! 4.0 Nemzetközi Licenc feltételeinek megfelelően használható.



**NEMZETI
KÖZSZOLGÁLATI
EGYETEM**
LUDOVIKA

TARTALOMJEGYZÉK

Etika és jog

- Mi történik, ha az emberek nem értik, hogyan működik a mesterséges intelligencia?
- Az algoritmikus bűnüldözés álomképe
- Elborítják a munkaadókat az MI-vel generált önéletrajzok

Trendek

- A jövő valósága, avagy hogyan alakítja át a világunkat az MI által vezérelt metaverzum?
- Felkészülés a mesterséges intelligencia eljövendő biológiai képességeire
- „Táv-írás” a mesterséges intelligencia korában

Működésben

- A futball ad újabb lökést Kína humanoid robotokkal kapcsolatos innovációinak
- Szerelem chatbotok idején: érzelmi kötődés az MI-modellekhez
- Utánaszámoltunk az MI energia-lábnyomának
- Úgy beszélsz, mint a ChatGPT





Etika és jog

Mi történik, ha az emberek nem értik, hogyan működik a mesterséges intelligencia?

Bármit is mondjanak a nagy technológiai vállalatok vezetői, a nagy nyelvi modellek nem tekinthetők intelligensnek az ember által ismert értelemben. Egy most megjelent cikk két frissen megjelent könyv tanulságait foglalja össze. Az első címe *Empire of AI: Dreams and Nightmares* in Sam Altman's OpenAI (Az MI birodalma: Álmodok és rémálmodok Sam Altman OpenAI-jában), a másik címe pedig *The AI Con: How to Fight Big Tech's Hype and Create the Future We Want* (Az MI-csalás: Hogyan küzdjünk a technológiai nagyvállalatok túlzásai ellen, és hogyan teremtsük meg a kívánt jövőt?). Mindkét munka – az előbbi implicit módon, az utóbbi explicit módon – azt sugallja, hogy az MI-iparág tevékenysége átveréseken alapszik. Ez nem azt jelenti, hogy a technológia nem figyelemre méltó, hogy nincs haszna, vagy hogy a megfelelő kezekben nem fogja átalakítani a világot. Azt jelenti, hogy a mesterséges intelligencia nem az, aminek a fejlesztői eladják: egy újfajta gondolkodó – és hamarosan érző – gép. Ezek az állítások egy fogalmi hibát tükröznek: a nagy nyelvi modellek nem „értene”k, nem tudnak és nem is fognak semmit sem „megérteni”. Nem rendelkeznek érzelmi intelligenciával és nem is okosak a szó semmilyen értelmes, vagy felismerhetően emberi értelmében. A nagy nyelvi modellek valószínűség-számítási alapon működő lenyűgöző eszközök, amelyekbe szinte az egész internetet betáplálták, és nem gondolkodás útján, hanem statisztikailag megalapozott becslések alapján állapítják meg, hogy melyik lexikai elem követi valószínűleg a másikat. A szerzők ismertetnek néhány nagyon zavarba ejtő esetet, amikor emberek tévképzeik megerősítéséért, vagy a tényleges emberi kapcsolatok pótlakaként fordulnak a gépi rendszerekhez. Bár számos MI-szolgáltatás valamilyen emberi tevékenységi területet hivatott felváltani, mindegyiküket az „antropomorfizálási hagyománynak” nevezett módszerrel reklámozzák: úgy beszélnek a nagy nyelvi rendszerekről (és ezen alapon is értékesítik őket), mintha azok emberi elmével rendelkeznének. Az elmúlt 30 év számos világot megváltoztató szilícium-völgyi technológiáját – elméletileg – az emberi boldogság, a kapcsolatteremtés és az önismeret növelésének eszközeként népszerűsítették, de ezek a gyakorlatban éppen az ellenkezőjét eredményezték. Érdekes adalék és kontraszt ehhez, hogy a legutolsó felmérések szerint az MI-szakértők 56 százaléka gondolja azt, hogy a mesterséges intelligencia jobbá teszi az Egyesült Államokat, de az amerikai nagyközönség ezt az optimizmust nem osztja, csak 17 százalékuk ért ezzel egyet.

[What Happens When People Don't Understand How AI Works](#)

Az algoritmikus bűnüldözés álmoképe

Több olyan téma is akad a mesterséges intelligencia technológiák fejlődésével, illetve különösen azok társadalmi beágyazódásával kapcsolatban, amelyek egyfajta tematikai örökzöldként újra és újra felbukkannak nem csupán a sajtóban, de a szakirodalomban is. Ezek között első helyen említhetjük a bűnüldözéshez, pontosabban a bűnmegelőzéshez, bűnüldözéshez, és igazságszolgáltatáshoz (tehát a társadalmi rend megőrzésének teljes vertikumához) kapcsolódó témákat. Ezen belül is az egyik téma az MI-vel támogatott, vezérelt megfigyelő rendszerek fejlesztése, telepítése és működtetése. A másik nagy téma, amely talán kicsit ritkábban bukkan fel, ám nagyságrendekkel nyugtalanítóbb perspektívákat vet fel, mint a surveillance-technikák: a csúcstechnológiák kombinációjával, elsősorban a mesterséges intelligenciával megvalósított technológia alapú előrejelző bűnmegelőzés. Sokféle elnevezése van, az elterjedt „predictive policing”-től kezdve, de a koncepció és technológia lényege: az új eszközrendszer segítségével jó előre észlelni és azonosítani a törvénybe ütköző cselekedetek lehetőségét, illetve azok potenciális elkövetőit. Első pillantásra kevés ennél pozitívabb alkalmazása lehet az új technológiáknak: megelőzni a bűnt, a bűnt. Sajnos azonban az alkalmazás elméleti alapjai (pl. az emberi lélektan) többnyire túlságosan is ingatag lábakon állnak. Így ezek a szándékuk szerint társadalmilag jótékony új eszközök a társadalmi diszkrimináció, igazságtalanság, jogtalanság előmozdítói lehetnek. A problémát ismét a közvélemény elé tuszkolta a brit kormány által most bevezetni kívánt új, MI-alapú bűn-előrejelző technológia, amely személyes egészségügyi és más adatok alapján minősíthetne embereket – bűnmegelőzési szempontból – "magas kockázatú" egyedeknek. Eközben Argentína új Mesterséges Intelligencia Biztonsági Egysége (Artificial Intelligence Unit for Security) gépi tanulást kíván alkalmazni a bűncselekmények előrejelzésére és a valós idejű megfigyelésre. Kanadában pedig olyan városok rendőrségei, mint Toronto és Vancouver, a predictive policing jegyében már most olyan eszközöket alkalmaznak, mint a Clearview AI arcfelismerő rendszer. Egyes amerikai városokban pedig a mesterséges intelligencia alapú arcfelismerő rendszert utcai térfigyelő kamerákkal kombinálják a gyanúsítottak nyomon követése érdekében.

[When it comes to crime, you can't algorithm your way to safety](#)

Elborítják a munkaadókat az MI-vel generált önéletrajzok

A mesterséges intelligencia rendszerek tömeges használata rengeteg előre nem látott problémát hoz létre – és ez fokozattan így lesz a jövőben is. A munkaügyi felvételi folyamat újabban tapasztalt problémái ennek az egyik megnyilvánulása. Az álláshirdetésre érkező válaszok újabban megnőtt tömege azt jelzi, hogy a jelentkezők kiterjedten alkalmazzák az MI-rendszereket a jelentkezések megírásában és azoknak a vállalati hirdetéseknek a kiválasztásában, amelyekre válaszolnak. Például egy egyszerű utasítással a ChatGPT beilleszti az állásleírásban szereplő összes kulcsszót a jelentkező önéletrajzába. Olyan MI-ágensek is vannak már, amelyek fizetett szolgáltatásként önállóan keresnek álláshirdetéseket a megbízóik számára, és önállóan jelentkeznek is a megbízóik nevében. A munkaerő toborzásával foglalkozó szakemberek szerint egyre nehezebb megmondani, ki az, aki valóban alkalmas és sok önéletrajz gyanúsán hasonló. Özönvízszerűen áradnak a jelentkezések. A segítséget „természetesen” az jelenti, hogy a felvételi jelentkezések elbírálásában, a kiválasztási folyamatban is aktívan alkalmazzák az MI-eszközöket. A LinkedIn, amely platform az állások meghirdetésének egyik legfontosabb színtere, újabban egy olyan MI-ágent kínál a hirdetőknak, amely képes önállóan

válaszúzeneteket írni, szűrőbeszélgetéseket folytatni a jelöltekkel, javaslatot tenni a legjobbnak tartott jelöltekre és természetes nyelvi eszközöket használva keresni potenciális alkalmazottakat. Ez a folyamat – a jelentkezés és az elbírálás automatizálása – azzal a veszéllyel jár, hogy a bizalom megszűnik. Nem lehet tudni, vagy csak nagyon nehezen kideríthető, hogy valójában ki is jelentkezik, mi a tényleges szakértelme, és az is elhomályosul, hogy a jelentkezés tényleges értékelése hogyan történik.

[Employers Are Buried in AI Generated Resumés](#)





Trendek

A jövő valósága, avagy hogyan alakítja át a világunkat az MI által vezérelt metaverzum?

A metaverzum - a különféle mesterséges valóságok összessége - néhány éve a mesterséges intelligenciát követő újabb hatalmas technológia-üzleti forradalomnak tűnt. Mint ismeretes a Meta milliárdokat fordított az ezzel kapcsolatos fejlesztésekre. Az újabb technológiai-gazdasági fejlemények nyomán kitűnt, hogy ez egy megalapozatlan elképzelés volt. Az azonban kérdés, hogy alapvető problémák okozták a kudarcát, vagy csak túl korán fogalmazódott meg? Frissen megjelent könyvében két neves szakember, a Grayling-Rosenberg szerzőpáros, töretlen optimizmusáról tanúságot téve azt állítja, hogy a metaverzum már a közeli jövőben megváltoztatja világunkat. Ez önmagában nem lenne új dolog, mert a technológia szüntelenül megváltoztatja világunkat, és különösen így volt ez az elmúlt száz évben. A szerzők szerint azonban most már a valóság, a realitás megváltoztatása van soron, mégpedig a következő évtizedben. A kissé merész állítást – mint említettük - két ismert szakember jegyzi, de az egyes részekhez is mértékadó szakmabeliek járultak hozzá, és az előszót Neal Stephenson (a metaverzum fogalmának megalkotója), az ismert sci-fi író és technológiai szakember írta. A könyv szerint egy új technológiai korszakba lépünk, amelyben a mesterséges intelligencia és az immerzív média minden szinten átalakítja a társadalmat, megváltoztatva azt, amit látunk, hallunk és tapasztalunk, így befolyásolva életünket. Az immerzív szemüvegek és az interaktív mesterséges intelligencia együttműködésének köszönhetően világunkban a valóság és a virtuális világ, az emberi és a mesterséges között húzódó határok gyorsan elmosódnak majd. A könyv ezt az átalakulást számos szempontból tárgyalja (egészen a geopolitikai következményekig). Úgy vélik, hogy számos adatvédelmi aggályra kell megoldást találni, szigorú és átgondolt szabályozásokra lesz szükség, hogy a jövő metaverzuma biztonságos és védett hely maradjon az emberiség számára.

[Our Next Reality](#)

Felkészülés a mesterséges intelligencia eljövendő biológiai képességeire

Az élettudományok az a terület, ahol különösen sokat várnak a mesterséges intelligencia alkalmazásától. A nagyon nagy lehetőségek azonban együtt járnak a nagyon nagy veszélyekkel is. Az OpenAI blogja azt ismerteti, hogy mit lehet tenni a veszélyek minimalizálására, a katasztrofális következmények kialakulásának megelőzésére e téren. Az igazán nagy problémát az okozza, hogy az élettudományi kutatások jelentős része potenciálisan kettős felhasználású. Az MI-rendszerek tudományos kutatást elősegítő képességei – mint például a biológiai adatokkal kapcsolatos következtetések, a kémiai reakciók előrejelzése vagy a laboratóriumi kísérletek irányítása – arra is megteremthetik a lehetőséget, hogy minimális szakértelemmel rendelkező embereket segítsenek biológiai fenyegetések létrehozásában vagy magasan képzett szakemberek számára könnyítsék meg a biofegyverek előállítását. Az OpenAI általános megközelítése az aktív megelőzésre összpontosít: nem szabad megvárni, hogy bekövetkezzék a biológiai fenyegetés. Ezért a modelljeik képzését úgy alakítják, hogy azokat a kéréseket, amelyek károkozáshoz vezethetnek, a modellek úgy válaszolják meg, hogy magas szintű ismereteket nyújtsanak, miközben elegendő részletet visszatartanak ahhoz, hogy megakadályozzák a nem szakemberek általi visszaéléseket. A modelljeikbe megbízható, a rendszerek teljes működését átfogó ellenőrzési funkciók vannak beépítve, hogy időben észleljék a kockázatos vagy gyanús, biológiával kapcsolatos tevékenységeket. Ha a szűrők alapján valami nem tűnik biztonságosnak, a modell választ blokkolják. Ez egyúttal egy automatikus vizsgálatot is elindít, amelybe szükség esetén emberek is bekapcsolódnak. A védekezésnek része a használatot szabályozó jogi és üzleti feltételek alkalmazása. A hibák belső kimutatásán több független munkacsoport (red teaming) is dolgozik, amelyek feladata a biztonsági intézkedések teljes körű áttörésének megkísérlése. Végül az egész modellt védő intézkedések összessége egy átfogó, összefüggő rendszert alkot, amely a fejlesztés és működtetés minden részletére kiterjed. Abból kiindulva, hogy nem minden szervezet működik a szükséges elővigyázatos elvek alkalmazásával, az OpenAI kezdeményező szerepet tölt be egy olyan magáncégek és állami szereplők által alkotott védelmi együttműködés kialakításában, amely a potenciális biológiai veszélyek semlegesítéséért dolgozik.

[Preparing for future AI capabilities in biology](#)

„Táv-írás” a mesterséges intelligencia korában

A generatív mesterséges intelligencia technológiák, ezen belül a nagy nyelvi modellek villámgyors térhódításával, szédületes fejlődésével és sokrétű adaptációjával párhuzamosan egyre-másra kongatják meg különböző szakmák képviselői a vészharangot, mondván: a technológia a hagyományos emberi tevékenységek megannyi területét teheti tönkre. Az ezzel ellentétes hangok, amelyek azt hangsúlyozzák, hogy az új technológia nem leváltja, hanem inkább segíti, megacélozza az embert számos működési terepén és területén, egyelőre kisebbségben vannak ebben a pánik keltette hangzavarban. Egy a rangos Springer kiadónál most megjelent tanulmány ezeket a pozitív kicsengésű, meglehetősen ritka „színfoltokat” gyarapítja. A tanulmány a még szokatlanul csengő nevű „táv-írást” járja körül, amelynek természetesen nincsen köze a XX. analóg távíró technológiájához. A „táv-olvasás” (distant reading) kifejezésre alapozva (amely nagy szöveges korpuszok számítógépes elemzését takarja), a „táv-írás” (distant writing) egy olyan újszerű jelenségként értelmezhető, ahol az irodalmi alkotások szerzői nagy nyelvi modelleket használnak narratívák előállítására, ám e közben megőrzik az

alkotó folyamat feletti kreatív kontrollt. A tanulmány (illetve az ott ismertetett distant writing koncepció) azt hangsúlyozza, hogy a nagy nyelvi modellek megjelenésével a szépírás evolúciója egy új állomáshoz érkezik (érkezhet) el. Lényege: az LLM-használat nem leváltja, hanem ellenkezőleg, kiteljesíti az emberi kreativitást. Az elmélet – noha még viszonylag sok kérdést hagy nyitva – mindenképpen figyelmet érdemel, amennyiben immár elméletileg megalapozott módon vetíti előre azt, hogy a generatív MI-technológiák hogyan alakítják át a közeljövő irodalmi alkotási folyamatait, és ezzel az egész irodalmi panorámát is.

[Distant Writing: Literary Production in the Age of Artificial Intelligence](#)





Működésben

A futball ad újabb lökést Kína humanoid robotokkal kapcsolatos innovációinak

Június hónap két utolsó napja egy különös nemzetközi „sportesemény” első rangadójától volt hangos. Az e-sport persze már jó ideje nem számít ténylegesen újdonságnak, de ez a rendezvény mégis eredeti versenyszámokkal hívta fel magára a figyelmet – és messze nem csupán a sportrajongók körében. A 2025 World Humanoid Robot Games, tehát a 2025. évi Humanoid Robot Világtorna nyitó rendezvénysorozatát ugyanis a RoBoLeague különleges labdarúgó mérkőzése jelentette. A meccs, amelynek során 3 humanoid robot játszott az ellenfél csapat hasonlóan humanoidokból kiállított gépi hármásával, abszolút elsőnek számított a világon: míg ugyanis korábban már játszottak egymás ellen távolról (valóságos emberek által) vezérelt robotsportolók, ennek a mérkőzésnek a szereplői közvetlen emberi beavatkozás nélkül, csakis a mesterséges intelligenciájukra támaszkodva küzdöttek az ellenfél hasonlóan MI-vezérelte csatáraival szemben. A gépek valós idejű döntéshozatalra, a játékostársakkal összehangolt mozgásra voltak képesek, ahogy természetesen emberi beavatkozás nélkül „szedték össze magukat” egy-egy esés után is. A játékokban résztvevő robotokhoz némileg hozzáigazított szabályoknak is köszönhetően a robotfocisták az emberekre jellemző gördülékenységgel, folyamatos mozgással cikáztak a pályán. A Pekingben megrendezett viadalon a főváros Információtudományi Egyetemének, a neves Csinghua Egyetemnek, és a Kínai Mezőgazdasági Egyetemnek a robotcsapatai küzdöttek fel magukat az élbolyba. Szakértők szerint a rendezvény további lendületet ad majd a kínai humanoid robot-fejlesztésnek, amely egyébként is szédületes ívet futott be az elmúlt néhány évben. A robotikának ez a különleges ága Kínában az elméleti kutatásoktól szinte hónapok alatt küzdötte fel magát az ipari megvalósítás szintjére. A szervezők úgy vélik, hogy ez a rangos és látványos esemény (amelynek sok szempontból megágyazott az áprilisban, szintén Pekingben megrendezett első humanoid robot félmaraton futóverseny) jó felvezetője lesz majd a nyár végén a kínai fővárosban helyet kapó Robotikai Világkonferencia szakmai programjainak is.

[Robotic football kicks off China's leap in humanoid robot innovation](#)

Szerelem chatbotok idején: érzelmi kötődés az MI-modellekhez

A köznapis világunkba mindennapi használati eszközként belépő nagy nyelvi modellek többen jelentenek a felhasználók jókora részének egy egyszerű „okos szoftverét”. A felhasználók meglepően nagy hányada érzelmi kapcsolatot épít ki az általuk használt chatbotokkal. Ennek természetesen megannyi fokozata lehet, az inkább csak szórakozásnak szánt játékos csevegéstől az egészen komoly érzelmi kötődésig, bizonyos értelemben a „szerelemig”. Ez utóbbi nyilvánvalóan éppen úgy plátói, mint Garcia Marquez regényalakjainak testiség nélküli emocionális hullámváza, de mégis, valóságos emberi érzelmek. A jelenségnek sokan nyugtalanító jelentőséget tulajdonítanak, míg más kutatók szemében ez az ember alapvető kötődés-igényének természetes megnyilvánulása csupán. A mindenképpen érdekes és szembeötlő témáról való társadalmi és szakmai diskurzust egyértelműen segíti majd az a kutatás, amelyet a közelmúltban tettek közzé a japán Waseda Egyetem kutatói. A szakemberek ugyanis egy mérőeszközt, egy skálát dolgoztak ki annak mérésére, hogy valójában milyen érzelmi szálak és hogyan alakulnak ki az emberi felhasználó és az általa használt MI-chatbot között. Az EHARS (Experiences in Human-AI Relationships Scale, azaz az Ember és MI közötti Kapcsolatok Skálája) felhasználásával végzett kísérleteik azt mutatták, hogy a nagy nyelvi modellek használata nyomán az emberek 75%-a volt hajlamos arra, hogy érzelmi kérdésekben tanácsadóként, lelki támaszként tekintsen az LLM-re, míg 39%-uk az MI-t állandó, megbízható szereplőnek tekintette az életében. A tanulmány 242 kínai résztvevőt vizsgált, közülük 108-an (25 férfi és 83 nő) teljesítették a teljes EHARS felmérést. A tanulmány két, az emberi kapcsolatokat tükröző MI-hez fűződő kötődési mintát azonosított: a kötődési szorongást és a kötődés elkerülését. A kutatók hangsúlyozták, hogy eredményeik nem feltétlenül azt jelentik, hogy az emberek valódi érzelmi kötődést alakítanak ki az MI-rendszerek iránt, hanem inkább azt, hogy az emberi kapcsolatokra alkalmazott pszichológiai keretek az ember-MI interakciókra is alkalmazhatók. A kutatás vezetője arra figyelmeztetett, hogy az MI-platformok kihasználhatják a sebezhető felhasználók érzelmi kötődését pénzszerzésre, vagy akár még rosszabb célokra is.

[Love in the Time of Chatbots: 75% of Users Turn to AI for Emotional Advice, Study Finds](#)

Utánaszámoltunk az MI energia-lábnymának

A mesterséges intelligencia technológiákkal kapcsolatos szédületes versenyfutás közepette (párhuzamosan a szinte exponenciálisan növekvő kapacitásokkal) egyre gyakrabban kerül elő az új technológiák környezetvédelmi, fenntarthatósági következményeinek a témája. Néhány hónap óta szinte közhelyszerű a szakirodalomban az MI-technológiák fejlesztésének (a modellek tréningezésének), majd később folyamatos működtetésének hatalmas vízigénye, hogy a szükséges energiamennyiségről már ne is beszéljünk. Ehhez kapcsolódóan az elmúlt hetekben egyébként már olyan fejtegetések láttak napvilágot, amelyek nem a számítási kapacitásokban, vagy a rendelkezésre álló adattömegben jelölik meg az MI-fejlesztések (és így a globális MI-versenyfutás) igazi kulcsát, hanem a rendelkezésre álló energiát jelölik meg az MI-nagyhatalmi pozíció feltételeként. A káprázatos teljesítményekre képes és a társadalmi működéseket radikálisan megújítani ígérő MI-technológiák ökológiai lábnyma tehát már jó ideje a szakmai (és politikai) közbeszéd témája, ám mindeddig kevés kézzelfogható, tényszerű információ került napvilágra a modellek és az eszközök valóságos környezeti hatásairól. A témáról most a tekintélyes Massachusetts Institute of Technology ad átfogó körképet, plasztikusan bemutatva azt a folyamatot, ahogyan a mesterséges intelligencia körül sűrűsödő

technológiai szektor egyre jobban eltávolodik az emergens iparágak korábbi energiahatékonyságától. A tanulmány szerzői külön kiemelik azt, hogy az MI-technológiák energiafelhasználása az esetek jó részében egyszerűen átláthatatlan. A transzparencia-hiány abból fakad, hogy a nagy technológiai cégek vonakodnak nyilvánosságra hozni az energiafelhasználáshoz, környezetvédelemhez, fenntarthatósághoz kapcsolódó adataikat. Az MIT-kutatás külön érdekessége, hogy részletesen bemutatja az MI-használathoz (tehát pl. az egyes lekérdezésekhez, kép-, vagy szöveg generálásokhoz) kapcsolódó energiaigényeket. A szerzők különösen aggályosnak tartják, hogy a technológiai cégek (pl. a nagy adatközpontok működtetéséhez) kifejezetten szennyező energiaforrásokat használnak. Ráadásul, a környezetterheléssel kapcsolatos plusz költségeket áthárítják a fogyasztókra.

[We did the math on AI's energy footprint. Here's the story you haven't heard](#)

Úgy beszélsz, mint a ChatGPT

A nagy nyelvi rendszerek hatására változik az emberi beszéd, és vélhetően a gondolkodás is. Az utóbbi években nagyon markáns bizonyos szavak sokkal gyakoribb előfordulása. Ennek az az oka, hogy a chatbotokat részben nyugat-afrikai emberek alkalmazásával tanítják, kihasználva az ottaniak angol nyelvtudását, és egyúttal nagyon olcsó béreit. Így lett a „delve” szó, amelynek gyakori használata a Nigériában beszélt angolra jellemző, általánosan tapasztalt megnövekedett előfordulása (az MI-rendszerek által előállított szövegeken kívül) ennek az egész folyamatnak a jelképe. A kutatások szerint az MI-rendszerekkel folytatott kommunikáció hatására egyneműsödik, színtelenné válik az emberi kommunikáció. Eltűnnek az egyedi jellemzők, sajátosságok, tájszólások, nyelvi tévesztések – egyszóval azok a vonások, amelyek egyedivé, emberivé teszik az emberi kommunikációt. Amikor körülöttünk mindenki a gépileg előállított „helyes” angolt használja, akkor ez egyben nyomás alá helyez mindenkit. A nyelvi homogenitás a gondolkodásbeli homogenitás elősegítője. Egyben a bizalom is megrendül ebben a helyzetben, mert nem tudjuk, hogy egy emberrel kommunikálunk, vagy egy gépi rendszerrel. Viszont a jövőben azzal is számolhatunk, hogy a gépi kommunikáció fejlődése során egyre jobban az egyes felhasználók nyelvi sajátosságaihoz igazodnak az általuk használt rendszerek megnyilvánulásai. Az nem képezheti vita tárgyát, hogy a mesterséges intelligencia továbbra is formálja-e a beszédmódunkat, mert bizonyosan formálni fogja. A szerző optimista meggyőződése szerint azonban a jövő nincs előre meghatározva, a technológia fejlődhet úgy is, hogy jobban tükrözze az emberi sokszínűséget, ahelyett, hogy elnyomná azt.

[You Sound Like ChatGPT](#)

