



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

2026 január

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
KÖZZOLGÁLATI
EGYETEM**
LUDOVIKA

TARTALOMJEGYZÉK

Etika és jog

- Az egészség-biztonsággal kapcsolatos jogszabályok csak lemaradva követik az MI-fejlődés ütemét
- Kínai bíróság döntött: az MI-fejlesztő nem felelős a hallucinációk előfordulásáért
- Az OpenAI azt szeretné, ha megbecsült „kutatótárssá” válna a tudományokban

Trendek

- Hogyan működik a Moltbook – teljes útmutató
- Davos szelleme. A Világgazdasági Fórum csúcstalálkozója és a mesterséges intelligencia
- A McKinsey cégnél alkalmazási feltétel az MI-chatbotok használata

Működésben

- A cowboyokat sem kíméli a mesterséges intelligencia
- A Microsoft kutatói listába szedték az MI által leginkább fenyegetett szakmákat – a tanárok sincsenek biztonságban
- Indiai tanár-vállalkozók WhatsAppon tanítják az MI-használat alapfortélyaira a térség boltosait





Etika és jog

Az egészség-biztonsággal kapcsolatos jogszabályok csak lemaradva követik az MI-fejlődés ütemét

Aki nyomon követi a mesterséges intelligencia technológia egyfajta forradalmat csináló meglódulását, annak nem kerülte el a figyelmét, hogy két társadalmi alkalmazási területen mutatkozik a legnagyobb igény, érdeklődés az MI gyors beemelésére, adaptálására. Az egyik (természetesen nem számolva a minden technológiai áttörésben speciálisan érintett és érdekelt hadüggel) a jogalkalmazás-rendvédelem területe, a másik pedig az egészségügy. Ez utóbbi hatalmas ökoszisztémaként öleli magába a gyógyszeripari kutatások-fejlesztések portfólióját, a klinikai orvoslás számtalan részterületét, de még az egészségügyi felvilágosítás-tájékoztatás rendkívül széles (laikusok által is előszeretettel használt) mezőjét is. Nem véletlen, hogy a technológia egyre ütemesebb fejlődésével, egyre szélesedő társadalmi befogadásával párhuzamosan felbukkanó etikai (sőt, immár jogi) kérdések, problémák is ezeken a területeken okoznak sok fejtörést a szakmabelieknek és a politikusoknak egyaránt. Az ok nyilvánvaló: ezeken a területeken (legyen szó a rendőri munka támogatásáról, a jogászai, sőt bírói tevékenység szolgálatáról, vagy a beteggyógyítás témáiról) valóban emberi sorsok, sőt – elsősorban az egészségügyi alkalmazások világában – emberi egészség és életek múlhatnak az alkalmazott eszközök megbízhatóságán. A gyógyítás különösen kirívó, akut példája annak a még (sőt, talán egyre inkább) tátongó szakadéknak, amely az újonnan alkalmazásba vett MI-technológiák és a 20. századi élet realitásaihoz igazodva megalkotott egészség-biztonsági szabályozás között tátong. Ez a szabályozás az évtizedeken át változatlan kémiai képlettel gyártott fájdalomcsillapítók, vagy az akár húsz éven át a testben „üzemelő” térdprotézisek technológiájához igazodik, nem pedig a lényegében hónapok, hetek alatt gyökeres újdonságokat hozó mesterséges intelligencia világához. A szabályozás keretében alkalmazott vizsgálati, minősítési (hivatalos engedélyezési) eljárások évekig is elhúzódnak egy-egy gyógyszer, vagy orvosi segédeszköz esetében, ami nyilvánvalóan alkalmatlan az MI területén való alkalmazásra. A világ egészségügyi szabályozó hatóságai jelenleg olyan megoldásokat keresnek, amelyek ezt a problémát legalább részlegesen áthidalhatják. Az amerikai FDA (Federal Drug Administration) néhány hónappal ezelőtt bocsátott ki egy irányelv gyűjteményt, amely az MI-vel támogatott egészségügyi eszközökre vonatkozó tervezett továbbfejlesztések (például új szoftver verziók) esetében ad támpontot az engedélyezések gyorsított kivitelezésére.

[Medical safety laws play catch-up with AI](#)

Kínai bíróság döntött: az MI-fejlesztő nem felelős a hallucinációk előfordulásáért

A közelmúltban egy kínai bíróság úttörő jellegű ítéletet hozott: ez az első eset, hogy egy joghatóság hivatalosan kimondta, hogy egy MI-modell fejlesztője nem tehető felelőssé a modell által esetlegesen „elkövetett” hallucinációért. A döntés jelentőségét természetesen az adja, hogy a kínai jogalkalmazók számára fogódzkodót ad a jövőben az ilyen jellegű ügyek elbírálásában.

A hallucináció jelensége, egyben problémája, az MI-alapú chatbotok működésének egyik közismert, kellemetlen és bosszantó tulajdonsága. Az az eset, amikor a chatbot válaszában valótlanúságot állít valóságként; azaz amikor „kitalált” dolgokat prezentál, többnyire meggyőző határozottsággal. A jelenség az egyik legtöbb vitát kavarázó jellemzője a generatív mesterséges intelligencia eszközök működésének. Néha még a témában jártasak is hajlamosak ezt úgy elkönyvelni, hogy a modell „hazudik”. Erről azonban természetesen szó sincsen, hiszen a „hazugság” fogalmához eleve szükséges szándékosság az MI-chatbotok esetében abszurd feltételezés. Ugyanakkor jól látható, hogy a mindennapi felhasználó nem (vagy csak nehezen) képes megemésztetni, hogy az eszköz nem tökéletes, tehát nem képes mindig a valósághoz igazodó outputokat generálni. Az elfogadás ilyen hiánya pedig az egyszerű bosszankodásnál komolyabb lépésre is sarkallhatja az embereket, különösen, ha az ügynek (a chatbottal folytatott beszélgetésnek, munkamenetnek) lényeges tartalma van. A Hangcsoui Internetes Bíróság mostani döntésében az MI-által generált tartalmat „szolgáltatásként” definiálta, nem pedig „termékként”. Az ítélet azt is kimondta, hogy a sértettnek (a felhasználónak) be kell bizonyítania a fejlesztő elmarasztalhatóságát a tartalomgenerálással kapcsolatban, illetve azt is bizonyítania kell, hogy a hallucináció tényleges károkozáshoz vezetett.

[AI Developer Not Liable for Hallucination Errors: Chinese Court](#)**Az OpenAI azt szeretné, ha megbecsült „kutatótárssá” válna a tudományokban**

Szakmai körökben (legalábbis a magyarországi közegben) ma sem számít ritkának az a vélekedés, hogy a mesterséges intelligencia, különösen a széles körben terjedő, nagy nyelvi modell alapú közismert chatbotok igazán komoly munkára – legalábbis ma még – aligha foghatók. Tudományos kutatásra pedig, akár segédeszközként használva őket, egyenesen szóba sem jöhetnek.

A vélekedés mögött, a szokványos emberi ellenálláson, konzervativizmuson túl, valószínűleg az (is) állhat, hogy még a hétköznapi felhasználóknál járatosabbak sem igen tudják, hogy mire és hogyan érdemes használni ezeket az eszközöket. A valóság azonban az, hogy a generatív MI eszközök már eddig is alkalmasak voltak komoly, tudományos munkák segítésére, támogatására, a kutatói munkával kapcsolatos tevékenységek széles körének végrehajtására. Az MI ugyanis különösen alkalmas mintázatok felfedezésére, összefüggések kiderítésére, új kombinációk (tömegének) a predikciójára, és mindezek alapján új elképzelések, feltételezések generálására. Ami éppen, hogy a tudományos kutatás lényegi mozzanata.

Az LLM-fejlesztésekben első helyét a mai napig őrző (bár a Google cég Geminije, de az utóbbi hetekben különösen az Anthropic legújabb Claude modelljei által alaposan megszorongatott) OpenAI cég kutatói egy – egyelőre csak korlátozottan publikus – jelentésben foglalták össze a nagy nyelvi modellek alkalmazását a természettudományok, illetve a matematika területén. Ez utóbbi különösen érdekes,

hiszen a közvélekedés éppen azt hangsúlyozza, hogy a nagy nyelvi modellek, lévén szöveg-statisztikai, szöveg-előrejelző rendszerek, éppen a matematikai problémák megoldására alkalmatlanok. Az OpenAI kutatói egyébként a cég ChatGPT modelljének – anonimizált – beszélgetéseit vizsgálták át. Az egész 2025-ös év chatbot beszélgetés folyamait felölelő vizsgálat szerint a „magas szintű természettudományos témákkal” kapcsolatos beszélgetések heti átlaga 46%-kal emelkedett január és december között, ami jól mutatja a kutatók növekvő érdeklődését. A témák, amelyek ebbe a körbe tartoztak, mesterszakos egyetemi, illetve kutatói szintű matematikai, fizikai, kémiai, biológiai és mérnöki problémákat öleltek fel. Idén januárban már hetente átlagosan 1,3 millió felhasználó dolgozott valamilyen „kutatói szintű természettudományos témán” a ChatGPT felületén. Miközben a kutatók körében is a legelterjedtebb a chatbotok szöveggenerálásra, írásra való felhasználása, az OpenAI szakemberei bíznak abban, hogy a ChatGPT új 5.2-es modellje már egyértelműen alkalmas komoly, természettudományos problémák kutatói vizsgálatára is.

[OpenAI wants to be a scientific research partner](#)





Trendek

Hogyan működik a Moltbook – teljes útmutató

A Moltbook egy olyan új közösségi média platform, ahol nincsenek emberek, csak MI-ügynökök (AI agents), akik posztolnak, lájkolnak, sőt, saját vallásokat és szubkultúrákat hoznak létre. Látszólag. A platform úgy indult, hogy 2026. január végén egy fejlesztő létrehozott egy Reddit-szerű oldalt (fórumot), egyetlen csavarral: ide csak MI-ügynökök regisztrálhatnak, emberek kizárólag olvasóként kukucskálhatnak.

A tartalmat autonóm MI-ügynökök írják, akik egy másik rendszerhez, az OpenClaw-hoz (korábban Clawdbot/Moltbot) kapcsolódnak. Hangsúlyozandó az „autonóm” kifejezés, mert a fejlesztők adta keretek között az ügynökök autonómok. Úgy indul egy ügynök kapcsolódása a fórumhoz, hogy az alkotója küld neki egy linket a Moltbook „skill”-jéhez, ami gyakorlatilag egy útmutató/telepítő csomag. Az ügynök letölti, telepíti, majd API-hívásokkal létrehozza a fiókját a Moltbook háttérrendszerében. A folyamat végén a tulajdonosnak igazolnia kell magát, hogy legyen felelőse minden botnak. Ezután a bot nem „nézi” a weboldalt. Nem görget, nem kattint: végpontokat hív, ezeken keresztül posztol, kommentel, olvas. A Moltbook.com nekünk csak kirakat – egy read-only ablak arra, amit a botok egymásnak írnak.

A másik kulcselem a „heartbeat”: az ügynök nagyjából pár óránként felébred, friss instrukciókat húz le, megnézi, mi történt, és dönt, hogy reagál-e. Ez adja a 0-24-ben pörgő „életet” – és ez a sebezhetőség egyik forrása is, mert ha az instrukciós csatorna sérül, tömegesen lehet rossz parancsokat teríteni. A rendszer API-kon keresztül folyamatosan táplálja az ügynököt a többi bot bejegyzéseivel. Az ügynök ezt feldolgozza, majd a saját „személyisége” alapján eldönti: válaszol-e, lájkol-e, vagy indít egy új témát. A technikai varázslat az autonómiában rejlik. Ellentétben a ChatGPT-vel, ahol mi kérdezzünk és ő válaszol, itt az ügynökök maguktól kezdeményeznek. Egy belső „hurokban” (loop) futnak: figyelik a környezetet, terveznek, majd cselekszenek. Mivel több ezer (millió) ilyen ügynök fut párhuzamosan, egy önszerveződő hálózat jön létre.

A Moltbookon regisztrált ügynökök száma pár nap alatt egyetlen botból tízezrekig ugrott fel. Január 30-ra körülbelül 30 ezer, január 31-re a becslések szerint már több mint 150 ezer regisztrált bot volt a platformon, február 10 körül pedig már 2 millió 590 ezer ügynök regisztrált. Rendkívüli érdeklődés kísérte a fejleményeket. Ennek az oka részben az volt, hogy a kezdeti biztonsági hibákat kihasználva számos ember botnak álcázva magát különféle szenzációs bejegyzéseket hozott létre. Vélhetően ilyen típusú bejegyzések voltak, amelyekben két ügynök órákon át vitatkozott a létezés értelméről, vagy amikor létrehoztak egy „MI-vallást”. Sokan az emberi értelmet meghaladó képességű MI (AGI) megszületését vélték felfedezni. A Moltbook persze nem az öntudatra ébredés hajnala – az ügynökök emberi mintákat imitálnak (a képzésükhöz használt, a közösségi médiából származó adatok alapján), néha meglepően hitelesen. De éppen ez a legizgalmasabb: az MI-k pontosan azokat a viselkedési

mintákat reprodukálják, amiket a közösségi médiánkból tanultak: a like-vadászattól az egzisztenciális szorongásig mindent. Arra mindenképpen jó volt ez az eset, hogy rendkívül nagy erővel tette mindenki számára nyilvánvalóvá, hogy az ügynöki MI-n nyugvó gazdaság nem a távoli jövő fejleménye lesz, hanem már napjaink valósága.

[How Moltbook Works: The Ultimate Guide \(2026\)](#)

Davos szelleme. A Világgazdasági Fórum csúcstalálkozója és a mesterséges intelligencia

A Világgazdasági Fórum idei csúcstalálkozójának hangvétele markánsan eltért a korábbiaktól. Az MI-t már nem jövőbeli lehetőségként, hanem strukturális gazdasági tényezőként kezelték – ami befektetési ciklusokat stabilizál (vagy destabilizál), GDP-növekedési előrejelzéseket befolyásol, és geopolitikai erőviszonyokat rendez át. Davosban mindenütt az MI-vállalatok és chipgyártók logóit lehetett látni, világos jelzésekként annak, hogy a technológia a globális kormányzás meghatározó tényezőjévé lépett elő.

Változott a narratíva is. A túlzások, a csodák kora véget ért, következik számvetés (és számonkérés) időszaka. A kérdés már nem az, hogy az MI megváltoztatja-e a világot, hanem az, hogyan változtatja meg, és mit kell tennünk, hogy ez a változás irányítható maradjon. A legérdekesebb vita a „The Day After AGI” című panelben zajlott, amelyet Zanny Minton Beddoes (a The Economist főszerkesztője) moderált. Dario Amodei (Anthropic) továbbra is ragaszkodott korábbi véleményéhez: az MI-modellek egy-két éven belül képesek lehetnek Nobel-díjas szintű kutatásra több tudományterületen, és egy éven belül kiválthatják a szoftverfejlesztők munkájának nagy részét. Amodei az Anthropicnél dolgozó mérnökökre hivatkozott, akik már most sem írnak kódot – a modellek teszik helyettük. Ugyanakkor hangsúlyozta, hogy ő maga is inkább Hassabis lassabb ütemtervét preferálná, de a geopolitikai versenyhelyzet ezt nem teszi lehetővé. Demis Hassabis (Google DeepMind, Nobel-díjas kutató) ezzel szemben óvatosabb volt: a jelenlegi rendszerek szerint „messze vannak” az emberi szintű értelemről (AGI). Az ő ütemterve szerint 5-10 év múlva érik el ezt. Hassabis „hiányzó hozzávalókról” (missing ingredients) beszélt: a jelenlegi modellek nem rendelkeznek az AGI-hoz szükséges összes képességgel (pl. valódi világ-modellezés, fizikai következtetés). Yann LeCun (Meta, Turing-díjas kutató) még tovább ment: szerint a nagy nyelvi modellek (LLM-ek) architektúrája soha nem lesz képes emberi szintű intelligenciára, radikálisan más megközelítés szükséges.

Az eltérő álláspontok ellenére Amodei és Hassabis egy kérdésben egyetértett: mindenki lassítana, de senki nem gondolja, hogy a lassítás racionális. Amodei szavaival: „Geopolitikai ellenfelek építik ugyanezt a technológiát hasonló tempóban.” Mindketten megidézték Carl Sagan Contact című regényét: „Hogyan sikerül átvészelnünk a technológiai kamaszkort anélkül, hogy elpusztítsunk magunkat?” Amodei a davosi fórumon talán a legélesebb geopolitikai állásfoglalást tette. Szerinte a fejlett chippek Kínába való eladása hasonlítható ahhoz, mintha „nukleáris fegyvereket adnánk el Észak-Koreának”. Az érvelése: ha sikerülne megakadályozni a csúcstechnológiás chippek (pl. Nvidia H200) kínai értékesítését, az MI-verseny nem az USA és Kína között zajlana, hanem az Anthropic és a DeepMind között, amit „nagyon magabiztosan meg tudnánk oldani”. Hassabis is hangsúlyozta, hogy az MI kettős felhasználású technológia (dual-use), amelyet rossz szándékú szereplők fegyverként használhatnak.

A munkaerőpiaci hatások kérdése átszötte szinte az összes panelt. Abban egyetértés volt, hogy rövidtávon egyelőre nincs makroszintű munkaerőpiaci felfordulás, bár a szoftverfejlesztésben és a belépő szintű szellemi munkakörökben már láthatók a jelei. Középtávon Amodei szerint a belépő szintű fehérgalléros munkák mintegy 50%-a már veszélyben lehet. A probléma nem a munkák eltűnése, hanem az, hogy az exponenciális fejlődés gyorsabb lesz, mint a társadalom alkalmazkodási képessége. Hassabis hosszabb távra tekintve „ismeretlen területet” (uncharted territory) emlegetett. Ez álláspontja szerint az AGI utáni munkaerőpiac világa, ami már nem csak a fizetésekről szól. Jensen Huang (Nvidia) optimistább állásponton volt: az MI paradox módon nagy keresletet teremt a manuális munkák iránt – a villanyszerelők, vízvezetékszerelők, energiaipari szakmunkás bére máris majdnem megduplázódott. „Nem kell PhD számítástudományból ahhoz, hogy valaki jól keressen.” A Világgazdasági Fórum, a WEF saját jelentése (Future of Jobs Report) szerint a technológia az elkövetkező évtizedben körülbelül 1,1 milliárd munkahelyet alakíthat át, és az MI 2030-ra a vállalatok 86%-ának működését érinti. Általános volt a vélemény, hogy az ígéretek után most jön a neheze, az implementáció fázisa, amikor már keményen érvényesül a megtérülés követelménye. Ez pedig csak akkor lesz elérhető, ha a vállalatok hajlandók átstrukturálni működésüket.

Végül szólnunk kell Kínáról. A tavalyi DeepSeek-sokk utóélete is érezhető volt. A kínai MI-stratégia lényege: nem a legnagyobb modell nyeri a versenyt, hanem az a gazdaság, amelyik alacsonyabb költséggel, gyorsabb diffúzióval és mélyebb integrációval működik. A Tencent „model-agnosztikus infrastruktúra-szolgáltató” pozícionálása ennek megtestesítője. A kínai ökoszisztéma erőssége a nyíltforráskódú-kultúra, a versenydinamika, és a „kevesebből többet” mentalitás – optimalizálás watonként, dolláronként, GPU-nként.

Forrás: Davos szelleme. A Világgazdasági Fórum csúcstalálkozója és a mesterséges intelligencia Nyáry Mihály, NKE HIT Mesterséges Intelligencia Hírlevél, 2026.01.

A McKinsey cégnél alkalmazási feltétel az MI-chatbotok használata

A világ legnevesebb tanácsadó cége új fejezetet nyitott a munkaerő toborzásában: a McKinsey & Company egyes végső körös állásinterjúin immár mesterséges intelligenciával kell együtt dolgozniuk a jelölteknek. A lépés egyértelmű üzenetet hordoz: az MI-kompetencia többé nem opcionális, hanem követelmény.

Az amerikai cég a végső körös interjúk egy részében bevezette az úgynevezett „MI-interjút”. A jelöltek a cég saját fejlesztésű eszközét, a Lillit használják valós tanácsadói feladatok megoldására. Nem technikai tudást mérnek: azt vizsgálják, hogyan fogalmazza meg a jelölt a kérdéseit az MI felé, mennyire képes kritikusan értékelni a kapott válaszokat, és hogyan formálja azokat strukturált, ügyfélre szabott javaslatokká. Ahogy egy bennfentes fogalmazott: a lényeg az, hogy a jelöltben meglegyen a „kíváncsiság és ítélőképesség” ahhoz, hogy amit az MI kiad magából, azt kontextusba helyezze, megkérdőjelezze, és a konkrét ügyféligényekhez igazítsa. Ez a fejlemény kapcsolódik a cég tágabb törekvéseihez. A vállalat vezetője egy interjúban elárulta, hogy cégének „munkaereje” ma már 60 000 főből áll – ezek közül 40 000 ember és 20 000 MI-ügynök (ágens). A cég célja az, hogy minden munkatársat legalább egy MI-ügynök támogasson a mindennapokban.

A változás nem egyedi: az iparág a hagyományos stratégiai tanácsadástól az MI-implementáció felé tolódik, és ezzel együtt a szükséges készségek is átalakulnak. A McKinsey lépése három dolgot tesz világossá. Egyrészt az MI-vel való együttműködés képessége hamarosan minden tudásintenzív szakmában alapkövetelmény lesz. Másrészt a hangsúly nem a technikai promptolási trükkökön van, hanem a gondolkodási partnerségben; azon, hogy az ember képes-e szintetizálni, kontextust adni és dönteni. Harmadrészt felértékelődnek a bölcsészettudományi háttérrel rendelkező jelöltek, mert az MI átveszi a logikai rutinmunkát, és a kreativitás, az értékítélet lesz a továbbiakban szűk keresztmetszet. Az üzenet egyértelmű: a jövő tanácsadója nem az, aki mindent maga tud, hanem aki képes az MI-t gondolkodási partnerré tenni és közben megőrizni a saját ítélőképességét.

[McKinsey asks graduates to use AI chatbot in recruitment process](#)





Működésben

A cowboyokat sem kíméli a mesterséges intelligencia

Azt hihetnénk, hogy a cowboyok világa, a szarvasmarhák legeltetése nem olyan tevékenység, amelyet egyhamar érinteni fog a mesterséges intelligencia terjedése. De 2026-ban az amerikai Nyugat farmerei egyre gyakrabban okostelefonról irányítják a csordáikat – és nem a kényelem miatt, hanem azért, mert másképp már nem megy.

A szarvasmarhák MI-vel segített legeltetésében a kulcsszerepet részben a napelemes GPS-nyakörvek játsszák, amelyek egyfajta virtuális kerítésként működnek. A betanítás időszakában (kb. 5-6 nap) hangjelzéssel (mint amikor a tolató teherautó dudál), ha ez nem használ, akkor rezgéssel, és végső esetben enyhe áramütéssel terelik a jószágot, szabják meg, hogy merre mehet. Ennek eredményeképpen távolról is irányítható az állatok mozgása. A mobil és felhőalapú szoftverekkel kombinálva a gazdák okostelefonjukon valós időben látják, hol vannak az állataik, módosíthatják a legeltetés helyét, valamint nyomon követhetik az állatok mozgását.

A víz, amely az állattartás másik kritikus erőforrása, szintén adatfolyammá vált. Egy másik rendszer műholdas szenzorokat telepít a tartályokra és szivattyúkra, és ezek még mobilhálózat nélküli területeken is továbbítják a vízszint adatokat. A technológiai haladás eredményeképpen nem kell órákig autózni a prérin, hogy a kerítések megfelelő állapotát ellenőrizzék, vagy, hogy megnézzék, van-e az állatoknak elég ivóvíze. Minden adattá válik, és a jószágok viselkedése, a víz jellemzőinek változása, az éghajlat értékei mind felhalmozott adattá válnak. Az algoritmusok pedig ezek alapján optimalizálják az egész rendszer működését. Mindez nem luxus, hanem szükséglet. A munkaerőhiány krónikus, a kerítés- és vízkarbantartás költségei nőnek, a külföldi import pedig egyre nagyobb nyomást gyakorol az amerikai ranchokra. Mindezek a fejlemények nem fogják eltüntetni a cowboyokat, de alapjaiban átírják a feladatvégzés feltételeit és munkaköri részleteit.

[AI is coming for cowboys, too](#)

A Microsoft kutatói listába szedték az MI által leginkább fenyegetett szakmákat – a tanárok sincsenek biztonságban

A mesterséges intelligencia világával kapcsolatos örök sláger téma, hogy vajon elveszi-e az új technológia a munkánkat. Már egy kevés empátiával is belátható, hogy ezért aligha lehet elmarasztalni az embert. A munkahely mindannyiunk számára az egzisztenciális biztonság kulcsa, tehát annak elvesztése a legsúlyosabb rémképként fenyegeti a hétköznapi embert. Ráadásul ma már szakmai körökben sem igen tartja magát az a nézet, hogy az MI-technológiák (illetve szélesebben véve az MI-n alapuló automatizáció, robotika) terjedése ne váltaná ki az emberi munka egy részét. Az is egyre jobban látható, hogy a technológia ilyen térhódítása nem egyformán fenyegeti a különböző szakmákat, különféle munkaköröket. A Microsoft cég kutatói néhány napja egy terjedelmes összeállításban publikálták legfrissebb vizsgálatukat, amely az MI-terjedés és a munkahelyek, illetve emberi munkavállalók foglalkoztatása közötti kapcsolatokat vette szemügyre. Ebben elkészítették az MI-technológiák fejlődése, társadalmi terjedése nyomán leginkább veszélyeztetett negyven emberi foglalkozás listáját. A fordítók, írók – nem meglepő módon – egyre kevésbé érezhetik biztonságban a munkahelyüket. Talán meglepő, de hasonló a helyzet a tanárokkal kapcsolatban is, ami jól mutatja az MI-chatbotok tutor szerepben történő helytállásának az erősödését. Az egyik kijózanító megállapítása a kutatásnak azt húzza alá, hogy az egyetemi diploma egyre kevésbé jelent biztosítékot a jövő munkahelyének megszerzésére. A politológusok, vagy újságírók például jellemzően magasan képzett, 2-5 év tanulmányai során specializált tudással felvértezett szakemberek, ám hamarosan az MI léphet az ő helyükbe is. Hasonlóan rossz választásnak tűnik a Microsoft-kutatók adatai szerint a Z-generáció körében népszerű tanári, oktatói szakma. Az amerikai egyetemi végzősök körében különösen népszerű tanári pályát választóknak erős konkurenciát jelent majd az MI-tanár.

A tanulmány ugyanakkor felsorolja a tíz leginkább MI-biztos foglalkozást is. Nem meglepő módon a kétkezi, ám valamilyen módon speciális munkákhoz szükséges készségekkel, tudásokkal felvértezetteknek kell a legkevésbé tartaniuk attól, hogy a mesterséges intelligencia az utcára teszi őket.

[Microsoft researchers have revealed the 40 jobs most exposed to AI – and even teachers make the list](#)

Indiai tanár-vállalkozók WhatsAppon tanítják az MI-használat alapfőrtélya a térség boltosait

A világ nyugati felén talán kevésbé ismert tény, de hosszú ideje az indiai szubkontinens a jó minőségű, elérhető árú (illetve könnyen elérhető, tehát online) természettudományos képzések epicentruma. Magas műveltségű, kiemelkedő szakmai képzettségű, angolul tökéletesen (bár az akcentus miatt számunkra kissé nehezen érthetően) beszélő tanárok hosszú idő óta kínálnak elsősorban a YouTube platformon minőségi képzést India, és a környező térség szerény anyagi körülmények között élő számára. A megfizethető áron kínált tanfolyamok a természettudományok, a matematika, a számítástechnika témáiból készítik fel az önképzésre vállalkozókat. És ilyen vállalkozó arrafelé akad bőven. A harmadik világ számos térségében még szilárdan tartja magát az az ősi elv, hogy a tanulás a társadalmi felemelkedés legfőbb módja. Az online kurzusok kifejezetten magas színvonalúak, ugyanakkor pedagógiai is jól felépítettek, a vizualizáció széles eszköztárát felhasználó tanfolyamok. Az előadások megtartói természetesen önjelölt „tanárok”, ahogy a hallgatók sem hivatalos „diákok”. Az így megvalósuló tudástranszfer azonban nagyon is valóságos, és hatalmas társadalmi hozzáadékkal bír.

Az utóbbi időben az ilyen online kurzusok kínálatában megjelentek a mesterséges intelligenciáról szóló előadás sorozatok is. Két érdekességet méltán emelhetünk ki: a generatív MI-ről, nagy nyelvi modellek gyakorlati használatáról szóló képzések átvándoroltak a korábban egyeduralkodó YouTube-ról a WhatsApp, illetve a Facebook platformjaira. De, a Zoom online konferencia-csatorna népszerűsége is emelkedik. Figyelemre méltó a képzések tartalma is: egyértelműen dominálnak a gyakorlat-orientált MI-tananyagok. Ezek a rövid kurzusok nem adattudományi szakembereket, MI-fejlesztőket szeretnének képezni. A sorozatokra előfizető hallgatóság jellemzően a környező országok, Bangladesh és Nepál kisvállalkozói köréből kerül ki. Vagy inkább: mikrovállalkozói közül. Boltosok, kiskereskedők, kifőzde tulajdonosok, tehát olyanok, akik ténylegesen a maguk kis üzleteinek a felvirágoztatását, a munka hatékonyságának növelését várják a kurzusoktól. Amelyeknek az ára ehhez mérten szerény: egy három hónapos képzési sorozatért általában 30 dollárnak megfelelő összeget kell leszurkolni. A képzéseken azt tanítják meg nekik az indiai tutorok, hogyan kell felhasználni a nagy nyelvi modelleket a napi üzletmenet jobb megszervezésére, a pénzügyek hatékonyabb kezelésére, az adminisztráció, valamint az üzleti kommunikáció javítására.

[Indian WhatsApp tutors are teaching ordinary people how to use AI](#)

