



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

2026 augusztus

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

- Emberközpontú nagy nyelvi modellek
- A mesterséges intelligencia használat felerősítheti a digitális egyenlőtlenségeket a világban
- A gyerekek túltesznek a mesterséges intelligencián a tanulásban – és máig nem tudjuk, miért

Trendek

- A Wall Street tőzsdeügynökei ugyanúgy kereskednének az MI számítástechnikai kapacitásokkal, mint az olajjal
- Tömeges elbocsátásokról ugyan nincsen szó, de a fiatal pályakezdők helyzete tovább romlott az MI miatt
- Biológiai mesterséges intelligencia modellek: új paradigmák az élet nyelvének hasznosítására
- Egyre több fiatal amerikai fordul a mesterséges intelligencia ellen – vajon miért?
- Az MI képesség javulása lassan időszerűvé teszi a teljes közoktatási koncepció újragondolását

Működésben

- Az MI nem fogja felváltani a radiológusokat, de gyökeresen átalakítja a munkájukat
- Múzeumok és az MI: kritikus döntések
- A gazdaság pálcát tört az MI-vel generált tartalmak felett





Etika és jog

Emberközpontú nagy nyelvi modellek

A mesterséges intelligenciák társadalmi összefüggéseivel foglalkozó egyik legjelentősebb és legrangosabb kutatóműhely, az amerikai Stanford Egyetem HAI Human-Centered Artificial Intelligence intézete – nevéhez híven – úgy véli, hogy a nagy nyelvi modell forradalom tényleges eredményeinek megítélésénél a hangsúlyt át kell helyeznünk a technikai paraméterekről az általuk kiváltott emberi hatásokra, változásokra. Az ok egyértelmű: az újabb és újabb modellek, rendszerek teljesítményének összevetésére rutinszerűen használt technikai metrikák szinte semmit nem árulnak el arról, hogy vajon hasznára vannak-e ezek a fejlesztések (és ha igen, akkor hogyan és mennyire) az őket használó embereknek.

Az intézet jelentős kutatások eredményeit összefoglaló jelentése néhány sommásnak látszó, de lényegi ajánlást fogalmaz meg az iparág fejlesztői oldala számára:

- a kompetitív előny immár nem az újonnan kifejlesztett modell technikai paramétereiben, hanem az általa nyújtott felhasználói előnyben ölt testet;
- a technológiai összehasonlító adatok (benchmark) korántsem jelentenek tényleges üzleti előnyt;
- a legígéretesebb új fejlesztések nem az automatizálásra, hanem a gép–ember együttműködésre helyezik a hangsúlyt;
- az „interakció-alapú tervezés” kulcsfontosságú: jobb interfészek, eligazodást segítő kontextuális elemek, személyre szabhatóság, a munkafolyamat-integráció nagyobb előnyt jelent, mint az egyes modellek folyamatos fejlesztései.

[Human-Centered Large Language Models](#)

A mesterséges intelligencia használat felerősítheti a digitális egyenlőtlenségeket a világban

A technológia fejlődésével, legutóbb pedig az MI-alapú ágensek túlságosan is ámokfutásszerű lépéseivel egyre élénkebb a vita (jól érzékelhető módon mindhárom fontos szinten: a szakmai műhelyekben, a döntéshozók és politikusok köreiben, sőt hovatovább a köznap felhasználók világában is) arról, hogy vajon átok-e az emberiség számára az MI-forradalom, vagy éppen a megváltást és a jövőt jelenti. Érdekes módon szinte alig kerül terítékre ennek a dilemmának egy nagyon fontos vetülete. Amikor ugyanis a mesterséges intelligencia ilyen lehetséges konzekvenciáit latolgatjuk, hajlamosak vagyunk sommásan a „társadalmakat”, esetleg „az emberiséget” említeni a következmények elszenvedőjeként, vagy az előnyök majdani élvezőjeként. Eközben pedig mindenki előtt nyilvánvaló, hogy az egyébként számos értelemben „globalizált” (összekapcsolt, sőt sokszor univerzalizált) földgolyót valójában nagyon markáns, és nagyon mély szakadék(ok) osztják fel. Amikor a gazdaságról vagy különösen manapság a nagypolitikáról gondolkodunk, látjuk a világ egy fejlett tömbjét különállni, sőt szembefeszülni egy másik, nagyobb blokkal, a manapság éppen Globális Délnek nevezett halmazal. Érdekes, hogy amíg a számítástechnika, a mobil kommunikáció gyors fejlődése mentén magától értetődően megjelent a „digitális szakadék” fogalmi eleme, addig szinte egyáltalán nem bukkan fel a prognózisokban az „MI-szakadék” nagyon is valóságos koncepciója. Pedig az elmúlt évtizedek harmadik világbeli digitalizációját elősegítő nemzetközi programokon dolgozó szakemberek mindannyian tudják, hogy minden újabb „transzformatív” technológia egy már rétegzett, tagolt, az intézményi kapcsolódásoktól a készségeken át az ipari háttérig számtalan formában „előkészített” társadalmi közegben telepszik meg.

Az is jól látható, hogy a mesterséges intelligencia fejlődése, elterjedése szempontjából meghatározó tényezők, különösen az általunk is külön tárgyalt MI-számítási kapacitások nagyobb tömbökben, klaszterekben alakulnak ki a világban. Nem meglepő, hogy a jelentősebb felmérések (pl. a Stanford AI Index Reportja vagy a Világbank adatsorai) egyértelműen azt mutatják, hogy a számítástechnikában eddig is meghatározó Egyesült Államok jelenti az MI-fejlesztés fizikai infrastruktúrájának (különösen az adatközpontoknak) a szupertársát. Az MI-fejlődés azonban nem csupán a fizikai tényezők függvénye; a speciális tudások, képzettségek, a sokat emlegetett MI-szaktudás szempontjából pedig – az OECD friss adatai szerint is – még markánsabb a néhány fejlett és a világ többi része közötti szakadék. Miközben az infrastrukturális és képzettségi területeken nem túl biztató a globális egyenlőtlenségek sokasodása, a legfontosabb terület nem is ez. Az, hogy ki hogyan részesül a korszerű technológiák, az MI-transzformáció áldásaiból (vagy átkaiból), az döntően azon áll vagy bukik, hogy ki az, aki globális szinten kialakítja és ellenőrzi a játékszabályokat. Márpedig, vallják a témát vizsgáló szakemberek, a globális MI-szabályozás és az MI-governance az a terület, ahol a Globális Dél országai végképp nem kapnak helyet az asztalnál, ahol ezek a döntések megszületnek.

[AI is Hyper-Scaling Digital Inequality](#)



A gyerekek túltesznek a mesterséges intelligencián a tanulásban – és máig nem tudjuk, miért

A nagy nyelvi modellek ma már rendkívül magas szinten használják az emberi nyelvet, de ehhez nagyságrendekkel több tanítóadatra van szükségük, mint egy gyermeknek az anyanyelve elsajátításához. Ezt a különbséget adathatékonysági szakadéknak nevezik. Míg egy modern LLM előtanítása több billió token feldolgozását igényelheti, egy kisgyermek már néhány tízmillió hallott szó után képes nyelvtanilag helyes mondatokat alkotni. Ahogy az egyik kutató megjegyezte, ha 30 millió szón tanítasz egy nagy nyelvi modellt, akkor egy értelmetlenségeket gyártó generátort kapsz, nem egy gyereket.

A különbség egyik legfontosabb magyarázata az lehet, hogy a gyermekek tanulása nem passzív adatbefogadás. A gyerekek aktívan alakítják saját tanulási környezetüket: mozognak, tárgyakat vesznek kézbe, kérdeznek, kísérleteznek, figyelik cselekvéseik következményeit, és új tapasztalatokat keresnek ott, ahol bizonytalanok. Vagyis nem egyszerűen „megkapják” az adatokat, hanem részben maguk választják ki, milyen információval találkozhatnak. Tanulásukat a látás, a hallás, a testi tapasztalat és a társas interakció egyszerre szervezi. Ez alapvetően eltér a mai nyelvi modellek működésétől, amelyek jellemzően előre összeállított adathalmazokon, passzívan tanulnak. A gyermekek ráadásul érzékenyek arra is, hogy ki közöl velük információt, mit tud az illető, és milyen szándékkal tanítja őket. Így nemcsak nyelvi mintázatokat sajátítanak el, hanem folyamatosan következtetnek a környezet, más emberek és saját tudáshiányaik szerkezetére is.

A kutatók megpróbálják a gyerekek tanulásához hasonló viszonyok között működő nyelvi modellek, illetve a gyerekek fizikai világát és interakcióit rögzítő fejkamerás felvételek segítségével vizsgálni azt, hogy a gyerekekre jellemző tanulási sajátosságok közül melyeket lehet gépi modellekben reprodukálni. Ennek keretében zajlik a BabyLM kutatási kezdeményezés és éves verseny, amelynek során megpróbálják a nyelvi modelleket emberi léptékű, viszonylag kis adatmennyiségen tanítani. A cél annak feltárása, hogy a gyermekekhez hasonló adathatékonyság milyen tanulási elvekkel és modellarchitektúrákkal közelíthető meg. Az eredmények arra utalnak, hogy kis adatmennyiségből is kialakíthatók meglepően erős nyelvi képességek, de a multimodális és aktív tanulás gépi megfelelői még kezdetlegesek. Az adathatékonysági szakadék csökkentése nem pusztán kisebb modelleket vagy kevesebb tanítóadatot jelenthet, hanem új tanulási paradigmát is: olyan MI-rendszereket, amelyek célzottan keresnek információkat, kísérleteznek, visszacsatolást használnak, és aktívan szervezik saját tanulási folyamatukat. Ez egyaránt vezethet hatékonyabb MI-hez és az emberi nyelvtanulás mélyebb megértéséhez.

[Kids outlearn AI – and we still don't know why](#)





Trendek

A Wall Street tőzsdeügynökei ugyanúgy kereskednének az MI számítástechnikai kapacitásokkal, mint az olajjal

Mostanában erősödik a szakmai vita arról, hogy vajon tényleg az adat (még hozzá természetesen a big data) tekinthető-e a digitális korszak „olajának”, vagy pedig az MI-technológiák fékezhetetlen étvágát kielégíteni hivatott számítási kapacitások (csúf és elferdített angol kifejezéssel a „compute”) jelentik ezt az abszolút kulcstényezőt. A vita most abban az értelemben eldőlni látszik, hogy a nagyon jelentős amerikai kockázati tőkebefektetési alap, az Andreessen Horowitz egy startup vállalkozása 33 millió dolláros fejlesztési tőkével éppen ebbe az üzletágba vág bele. A cél egyértelműen az, hogy az MI-gazdaság kulcstényezőjeként érzékelt számítási kapacitásokkal árutőzsdei tételként kereskedhessenek már a közeli jövőben is. Miközben az MI-technológiák, a gyorsan kiépülő és szélesedő mesterségesintelligencia-alapú gazdaság roppant számítási igényeiről viszonylag sok szó esik, kevesebbszer bukkan fel ennek a tényezőnek az elképesztően nagy piaci értéke. A Goldman Sachs globális tanácsadó cég úgy prognosztizálja, hogy 2026 és 2031 között 7,6 billió (tehát amerikai mértékegységgel 7,6 „trillion”!) dollárt fognak Földünkön befektetni számítógépes kapacitások, nagy adatközpontok és a hozzájuk kapcsolódó energiaszolgáltatók kiépítésére. Ehhez a hatalmas „költséghez” azonban jelenleg még nem áll rendelkezésre a megfelelő pénzügyi-pénzpiaci infrastruktúra.

Szakemberek rávilágítanak arra is: a számítási kapacitások tőzsdére vitele számos új jelenséggel szembesíti majd a tőzsdei alkuszok világát. A számítási kapacitás ugyanis, szemben például a sokat emlegetett olajjal, természetesen nem megfogható, fizikailag nem érzékelhető árucikk. És ami még ennél is sajátosabb: nem is „állandó” egység abban az értelemben, ahogy egy hordó olaj mindig egy hordó olaj marad. Az újabb és újabb MI-chipek jelentős teljesítménynövekedést ígérnek, és ez jelentősen befolyásolja a régebbi mikrochipek piaci (árutőzsdei) értékét is.

[Wall Street wants to trade AI compute like oil. Should it?](#)

Tömeges elbocsátásokról ugyan nincsen szó, de a fiatal pályakezdők helyzete tovább romlott az MI miatt

A Stanford Digital Economy Lab frissítette az MI foglalkoztatásra gyakorolt hatásáról szóló tanulmányát. A tanulmány címe beszédes: „Kanárik a szénbányában? Hat tény a mesterséges intelligencia legújabb foglalkoztatási hatásairól” (Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence, 2025. augusztus). Az ADP cég (Automatic Data Processing, jelentős HR- és bérszámfejtő vállalat) adatai alapján vizsgálták, hogyan változott a foglalkoztatás a ChatGPT 2022-es megjelenése óta, különös tekintettel az életkorra és az MI-kitettségre. A kutatás nem talált bizonyítékot arra, hogy az MI eddig a gazdaság egészében széles körű munkahely-kiszorulást okozott volna, ugyanakkor markáns eltérés látható a fiatal munkavállalók körében.

Az MI-nek legerősebben kitett foglalkozásokban dolgozó 22–25 évesek foglalkoztatási szintje 2026 júniusára mintegy 19%-kal maradt el attól a szinttől, ahol akkor lenne, ha az MI-nek kevésbé kitett, hasonló korú munkavállalók foglalkoztatásával azonos ütemben alakult volna. Ez a különbség 2025 júliusában még 15% volt. A tapasztaltabb munkavállalóknál nem figyelhető meg hasonló lemaradás. Az alkalmazkodás elsősorban abban jelenik meg, hogy kevesebb fiatal vesznek fel, nem pedig a munkaviszonyok megszűnésének növekedésében. A visszaesés főként azokban a foglalkozásokban koncentrálódik, ahol az MI használata elsősorban emberi feladatok automatizálásához kapcsolódik. Ahol az MI a munkavállalók tevékenységét kiegészíti, a foglalkoztatás stabil vagy növekszik. A tanulmány új eredménye, hogy fontos különbséget azonosít a kodifikált és a hallgatólagos tudás között: a fiatalok foglalkoztatása inkább a formalizált, dokumentált tudásra épülő munkakörökben csökken, míg a tapasztalati, nehezebben reprodukálható tudásra támaszkodó foglalkozásokban a tapasztalt munkavállalók helyzete kedvezőbb.

A szerzők hangsúlyozzák, hogy ezek leíró mintázatok, nem oksági bizonyítékok. Bár több alternatív magyarázat – például a kamatkörnyezet, a távmunka vagy az ágazati összetétel – is lehetséges, ezek önmagukban nem magyarázzák meg az eltéréseket. Az értelmezést az sem segíti, hogy az eredmények érzékenyek bizonyos módszertani választásokra. A kutatók ezért nem tekintik lezártnak az MI munkaerőpiaci hatásairól szóló kérdésre adott választ, hanem folyamatos mérésre és havi frissítésre töreksenek.

[No Widespread Displacement, but the AI Employment Gap for Young Workers has Widened to 19%](#)

Biológiai mesterséges intelligencia modellek: új paradigmák az élet nyelvének hasznosítására

Az MI egyre mélyebben alakítja át a biológiai kutatást (is), és új lehetőségeket nyit többek között a genomikában, a proteomikában és a gyógyszerfejlesztésben. A biológiai adatokon tanított modellek gyorsabban fejlődnek, mint a rájuk vonatkozó szabályozási és irányítási keretek. Ez egyre sürgetőbb kérdéseket vet fel a tudományos megalapozottsággal, a gyakorlati alkalmazhatósággal és a megfelelő irányítási mechanizmusokkal kapcsolatban. Az Európai Bizottság kutatóközpontjának új kutatása 480 biológiai MI-modellt vizsgálva azt elemzi, mire képesek ezek a modellek, milyen infrastruktúrára van szükség a fejlesztésükhöz és használatukhoz, mennyire állnak készen a gyakorlati alkalmazásra, valamint milyen szakpolitikai kérdéseket vetnek fel az Európai Unió számára.

A fenti kutatás bemutatásakor a biológiára többszintű információs rendszerként is tekinthetünk. A DNS tárolja az információt, az RNS közvetíti és szabályozza, a fehérjék pedig működéssé alakítják. A sejtek mindezt sejtszintű funkciókba rendezik, a szövetek és az élő szervezetek szintjén pedig ebből alakulnak ki az élettani folyamatok. Minden szintnek megvan a maga „szókincse” – például a nukleotidok és az aminosavak ábécéje, a génexpressziós állapotok vagy a morfológiai jellemzők –, valamint a maga „nyelvtana”, vagyis az információ működését meghatározó szabályrendszere, például a génszabályozás, a fehérjehajtogatódás, a jelátvitel vagy a biológiai szerveződés. Ezek a szintek nem függetlenek egymástól: az információ folyamatosan áramlik közöttük, és ugyanaz a biológiai állapot egyszerre több szerveződési szinten is értelmezhető.

Ez a szemlélet, amelyet olykor az „élet nyelveként” (language of life) írnak le, új megvilágításba helyezi azt is, hogy valójában mit csinálnak a biológiában alkalmazott MI-modellek. Olyan rendszerekről van szó, amelyek megtanulják értelmezni a biológiai információ különböző szintjeit, kapcsolatot teremteni közöttük, illetve új biológiai mintázatokat és struktúrákat létrehozni. A fehérjenyelvi MI-modellek a szekvenciaadatokból az aminosavak „nyelvét” sajátítják el; a genomikai MI-alapmodellek a nukleotidszekvenciák szabályszerűségeit tanulják meg; az egysejtes MI-modellek a sejtek különböző állapotait, a patológiai MI-modellek pedig a szövetek morfológiai sajátosságait ragadják meg. A kutatás élvonalát egyre inkább azok a modellek jelentik, amelyek képesek egyszerre működni a fent említett szinteken, és a biológiát nem egymástól elkülönülő területek összességeként, hanem egységes, több léptéken szerveződő információs rendszerként kezelik. Éppen ez az elmozdulás teszi stratégiai jelentőségűvé a területet a gazdaság és a társadalom szempontjából: az élenjáró modellek esetében az MI-rendszerek a mai fejlettségük szintjén már nem csupán a biológiai kutatásokra ráépülő, szűk feladatokra alkalmazott előrejelző eszközök. Ezek a rendszerek már képesek arra, hogy a biológiai információt különböző szerveződési szinteken dolgozzák fel és kapcsolják össze. Ennek jelentőségét az uniós szakpolitika is felismerte, és az EU szakpolitikai keretei egyre hangsúlyosabban támogatják az ilyen irányú fejlesztéseket.

[Biological AI models: new paradigms to leverage the language of life](#)

Egyre több fiatal amerikai fordul a mesterséges intelligencia ellen – vajon miért?

Ha valaki úgy gondolná, hogy a fiatal generációk, a gyakran emlegetett „digitális bennszülöttek” kitörő örömmel fogadják a mesterséges intelligencia fejlődését és gyors terjedését – az bizony messze jár az igazságtól. Miközben a tizenévesek, huszonevesek egy rétege nyilvánvalóan aktív, kísérletező felhasználója az új technológiák lehetőségeinek, addig a korosztálybeli attitűdök nem ritkán az enyhe érdektelenségtől egészen a kifejezett ellenszenvig terjednek. A jelenséggel széles körben találkozó felsőoktatási dolgozók, tanárok jó néhány meglepettől tudnának beszámolni, ám az ilyen, angol szóval „anekdotisztikus”, vagyis egyéni élményen alapuló tapasztalatok mellett még ritka az empirikus alátámasztás. Ilyen például a Gallup véleménykutató óriáscég egyik idei felmérése, amely – a 14 és 29 év közötti amerikai korosztályok MI-vel kapcsolatos attitűdjeit vizsgálva – arra az eredményre jutott, hogy a mesterséges intelligenciához pozitív várakozással viszonyulók aránya a korcsoporton belül egyetlen év alatt jelentősen (36%-ról 22%-ra) visszaesett; eközben a kifejezetten ellenséges érzületűek aránya folyamatosan emelkedett (22%-ról 31%-ra).

Az attitűdök részben közvetlenül a technológiához és annak megnyilvánulásaihoz kapcsolódnak, de ugyanakkor jelentős az új technológiák szélesebben vett társadalmi, gazdasági és ökológiai következményei nyomán kialakuló pozitív vagy negatív érzület. Sok válaszadó fiatal érzékeli az online platformokon megjelenő (és tömegesen az MI segítségével előállított) tartalmak rohamos uniformizálódását, elsekélyesedését. A Harvard Egyetem kutatói által végzett vizsgálat ugyanakkor azt is rögzítette, hogy az ún. Z-korcsoporthoz igen nagy arányban tart a mesterséges intelligenciák szellemi képességeket visszafejlesztő hatásától: az általuk megkérdezett minta 79%-a vélekedett úgy, hogy az MI-használat elkenyelmeseíti az embert; 62%-uk pedig egyenesen azt vélte, hogy az ilyen technológiák használata „elbutítja” a felhasználót. Természetesen – a korosztályra jellemző erős társadalmi érdeklődés következtében – az MI-technológiákkal kapcsolatos attitűdök alakulásában fokozott erővel jelennek meg a makrotársadalmi következményekkel (pl. munkahelyek elvesztése), illetve különösen a környezetvédelemmel kapcsolatos aggályok is.

[Why young Americans are turning against AI](#)

Az MI képesség javulása lassan időszerűvé teszi a teljes közoktatási koncepció újragondolását

A vita nem csitul, sőt új lendületet kap. „Befellegzett-e az oktatási rendszerünknek az MI térhódítása nyomán?” – a kérdéssel kapcsolatos legújabb álláspontot most nem kisebb intézmény, mint a világ egyes számú műszaki-természettudományos felsőoktatási és kutatóhelye, az amerikai MIT ad hoc bizottságának jelentése hozza el hírlevelünkbe. Az Eric Klopfer és Samuel Madden professzorok által vezetett testület arra vállalkozott, hogy körbejárja az MI-technológiák által jelentett kihívásokat; azokat, amelyekkel nem holnap, hanem már ma szembe kell néznie a minőségi felsőoktatásnak. Az eseti vizsgálóbizottság megállapításainak – riasztó horderejű – összegzése leszögezi, hogy az MI-rendszerek már ma is elfogadható módon („hitelesen”) képesek teljesíteni bármilyen, az egyetemi alapképzésben előforduló feladatot. Vizsgálataik szerint a Massachusetts Institute of Technology képzési tervében szereplő valamennyi esszéfeladatot, problémamegoldó munkát vagy programozási feladványt „elfogadható színvonalon, meggyőző érveléssel alátámasztva” képesek teljesíteni a mesterséges intelligencia eszközök.

Ezekre a megállapításokra a világ felsőoktatási intézményeiben dolgozó szakemberek egy jelentős része úgy tekint, mint valamiféle nagy tekintélyű megerősítésére annak, amit „már mindenki évek óta tud és vall”. Az MIT bizottságának jelentése azonban nem csupán a felsőoktatásban rohamosan terjedő „nem etikusan eszközhasználatra” (vagyis a mesterséges intelligencia eszközökkel támogatott csalásra) egyfajta szakmai megerősítés. Az ad hoc bizottság vizsgálata arra a következtetésre jutott, hogy a széleskörű MI-használat, a közvetlen etikai-jogi problémákon túlmenően nagyon súlyos közösségi következményeket is generál. Azt találták, hogy a chatbotok elterjedt használatával párhuzamosan kezdett leépülni az egyetemeken közösségi kultúrája is: egyre csökken a diákok részéről a személyes konzultáció (tehát az oktatói fogadóórák) igénybevétele; visszaszorul az online szakmai fórumokon való hallgatói részvétel; leépülőben van a személyes részvételen alapuló diák tanulócsoportok munkája a kollégiumokban, könyvtárakban és más közösségi terekben.

Összegezve: az MIT szakemberei arra jutottak, hogy az MI-kérdés már messze nem csupán „puskázási” kérdés. Vélhetően a teljes felsőoktatással kapcsolatos társadalmi elképzelésünket, az egyetemi oktatás

egész koncepcióját kell újragondolni, hogy a változó technológiai környezetben is megőrizzük ezek kulcsfontosságú társadalmi funkcióit.

[MIT Warns That AI Can Now Credibly Complete Pretty Much Any Undergrad Assignment, Considers Overhaul of Entire Educational Model](#)





Működésben

Az MI nem fogja felváltani a radiológusokat, de gyökeresen átalakítja a munkájukat

Geoffrey Hinton, aki később Nobel-díjat kapott, és az MI-keresztapjaként is ismert, 2016-ban azt jósolta, hogy öt éven belül számítógépek válthatják fel a radiológusokat. Ezt a jóslatát azóta is számtalanszor felidézik, meglehetősen élcélődve, hiszen ez nem következett be, sőt a szakma létszáma tovább nőtt. Ugyanakkor abban igaza volt, hogy az MI mára olyan „szilíciumalapú kollégává” vált, amely bizonyos diagnosztikai feladatokban eléri vagy meghaladja az emberi teljesítményt.

Hinton várakozásaival szemben a radiológia lett mára az egyik legjobb példa arra, hogy a mesterséges intelligencia nem feltétlenül váltja ki az emberi szakértőket, hanem alapvetően átalakítja munkájukat. A radiológia (orvosi képalkotás és a képi információ értelmezése) az egészségügyi MI egyik legfejlettebb alkalmazási területe. 2026 elejére az FDA (az Egyesült Államok Élelmiszer- és Gyógyszerügyi Hivatala) által engedélyezett mintegy 1400 MI-alapú orvostechikai eszköz háromnegyede radiológiai célokat szolgált. Ezek egy része gyorsítja az adminisztratív és értékelési munkát, más rendszerek pedig olyan eltéréseket is felismerhetnek, amelyeket az emberi szem nehezen vagy egyáltalán nem érzékel. Ez különösen fontos, mert a diagnosztikai képalkotás során előforduló emberi hibák aránya becslések szerint 3–5 százalék. A fő kérdés napjainkban nem úgy vetődik fel, hogy az ember vagy az MI „jobb-e”, hanem úgy, hogyan lehet a kétféle intelligencia erősségeit összehangolni. Az MI fáradhatatlanul képes képpontok és mintázatok elemzésére, míg a radiológus a betegségekről, a klinikai összefüggésekről és az adott páciensről rendelkezésre álló tudást viszi be az orvosi döntésbe. A nehézség abból fakad, hogy a neurális hálózatok gyakran „fekete dobozként” működnek, így az orvosnak úgy kell felügyelnie döntéseiket, hogy azok belső logikája nem feltétlenül átlátható. Ez új kognitív kockázatokat is teremt. A radiológus túlzottan megbízhat az MI-ben, kritikátlanul elfogadva hamis pozitív vagy hamis negatív eredményeit, de az ellenkező véglet is veszélyes: néhány látványos gépi hiba miatt indokolatlanul elutasíthatja az egész rendszert. Ezért a megfelelő képzés, az MI teljesítményének folyamatos nyomon követése és az egyes döntésekhez kapcsolódó megbízhatósági becslések kulcsfontosságúak. A cikk végkövetkeztetése szerint nem az MI fogja leváltani a radiológusokat, hanem az MI-t használó radiológusok kerülhetnek előnybe azokkal szemben, akik nem alkalmazzák ezeket az eszközöket.

[AI won't replace radiologists, but it will dramatically change their jobs](#)

Múzeumok és az MI: kritikus döntések

A múzeumok stratégiai és etikai döntések tömegével szembesülnek a mesterséges intelligencia gyors terjedése miatt. Az MI-vel kapcsolatos intézményi döntések szükségképpen ideiglenesek: a technológia, az árak, a szabályozás és a társadalmi megítélés olyan gyorsan változik, hogy folyamatos

nyomon követésre és rugalmas alkalmazkodásra van szükség. Két nagy döntési terület különíthető el. Az első az értékalapú kérdések köre. Ide tartozik az MI-eszközök „provenienciája”, vagyis, hogy milyen adatokon tanították őket, sérültek-e szerzői jogok az adatok felhasználása során, illetve mennyire egyeztethető össze használatuk a múzeum küldetésével. Szintén fontos a környezeti terhelés, az algoritmikus torzítások, valamint az MI foglalkoztatásra és az alkotói szakmákra gyakorolt hatása. Lehetséges megoldásként felmerül a kisebb, helyben futtatott, ellenőrizhető és jogtiszta betanított modellek alkalmazása, azaz olyan megoldásokra van szükség, amelyek esetében biztosítható a múzeumok sajátos értékeinek tiszteletben tartása.

A második terület a gyakorlati működésé. A múzeumoknak mérlegelniük kell saját technológiai és emberi kapacitásait, az MI-szolgáltatások gyors változásából eredő kockázatokat, a hosszú távú költségeket, valamint azt, hogy az MI hogyan hat a munkatársak autonómiájára, stresszére és munkahelyi elégedettségére. Kiemelt fontosságú kérdés a reputáció és a közbizalom biztosítása. A túlzott vagy átláthatatlan MI-használat ütközhet azzal az elvárással, hogy a múzeumok az emberi tudás és alkotás hiteles őrzői legyenek. Ezért az MI-stratégiát széles körű intézményi párbeszédre, világos értékekre és rendszeresen felülvizsgált döntési keretre kell építeni.

[Museums and AI: Critical decisions](#)

A gazdaság pálcát tört az MI-vel generált tartalmak felett

Az egyre sokrétűbb támogatásra alkalmas LLM-technológiákat elsősorban üzleti/munkafolyamatok támogatására veszik igénybe a nem specializált tudású, tehát programozáshoz nem értő tömegek. A felhasználók körében a második helyen továbbra is a szövegek íratása, a különféle tartalmak generáltatása áll, tehát az ilyen feladatokra előszeretettel veszik igénybe az MI képességeit, tudását. A szakirodalomban és a közbeszédben is ez főként etikai szemüvegen át vizsgálva merül fel, tehát az ilyen tartalmakkal kapcsolatban elsősorban az ma a kérdés: vajon jogilag (például a szerzői jog keretei között) vagy erkölcsileg – ez az egyébként nagyon széles felhasználói rétegeket lefedő felsőoktatási, közoktatási felhasználással kapcsolatos viták és dilemmák világa – megengedhető-e az ilyen LLM-mel generált tartalmak használata, közreadása. Van azonban ennek a kérdésnek egy merőben más, de nagyon is fontos vetülete: nevezetesen az, hogy vajon mit „ér” – a szónak most nem a megbecsültség, az erkölcsi elismertség szempontjából, hanem a lukratív hasznok oldaláról vett értelmében – az így előállított tartalom? Érdekes módon (érdekes azért, hiszen a modern társadalmakban a haszonelvűség meglehetősen erősen jelen levő szempont) ez mindeddig ritkán merült fel a vitákban, társadalmi egyeztetésekben. Pedig indokolt lenne, hiszen ma már a hétköznapi felhasználó is egyre inkább tapasztalja, érzi, hogy a különböző online felületeket, fórumokat szinte elborítják az MI-vel generált tartalmak. A mostanában készülő vizsgálatok, felmérések egyértelmű jelzést látszanak sugározni: a felhasználók csupán elhanyagolható mértékben hajlandók fizetni az ilyen géppel generált tartalmakért. Talán az sem véletlen, hogy ma már széles körben honosodik meg az ilyen írásművek, képi- és hanganyagok minősítő jelölésére az „AI slop”, azaz az „MI moslék” elnevezés. Az MI-vel készített tartalmak szempontjából tehát maga a piac látszik kimondani az (elmarasztaló) verdiktet. Érdekes ugyanakkor továbbgondolni ezt: a tömeges MI-szemét online megjelenésének negatív következményei mellett ugyanis egy kifejezetten pozitív hozadéka is van/lesz a tendenciának. Nevezetesen az, hogy az

MI-technológia alkalmazásának tömeges terjedése éppen hogy felértékeli a valós emberek által teremtett tartalmakat, legyen szó szövegről vagy képről.

[The Economy Has Spoken: Stuff That's AI-Generated has Almost Zero Value](#)

