



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

2026 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

- Mesterséges intelligencia Európáért, Európa saját nyelvein
- Elszabadulhatnak-e az MI robotok? – könnyebben, mint azt gondolnánk
- Az MI-ágensek gyakorlott kutatónak is beválnak - bár néha a kutatás manipulálásától sem riadnak vissza

Trendek

- Ez nem fog eltűnni: a pályakezdők állásválsága az MI miatt
- Az MI-gazdaság helyzete
- Kikerülnek az egyetemekről az „MI-bennszülött” diplomások első hullámai – és a munkaadók finoman szólva nem túl boldogok

Működésben

- El tudja-e végezni ezt a munkát egy gép? – az MI, a fogyasztóra hárított munka és az „önkiszolgáló gazdaság”
- Az ENSZ a mesterséges intelligencia növekvő környezeti „lábnyomára” figyelmeztet
- A mesterséges intelligencia már jobb diagnosztá, mint egy körzeti orvos – legalábbis egy szimulált vizsgálat szerint





Etika és jog

Mesterséges intelligencia Európáért, Európa saját nyelvein

A Microsoft Open Innovation Center vezetője úgy látja, hogy Európa digitális jövője szempontjából stratégiai kérdés, hogy a mesterséges intelligencia ne csak a nagy világnyelveken, különösen ne elsősorban angolul működjön jól, hanem Európa saját regionális és kisebbségi nyelvein is. Bár az Európai Unióban több mint hatvan ilyen nyelvet beszélnek, ezek a jelenleg fejlesztett MI-rendszerekben erősen alulreprezentáltak. Ennek oka főként az, hogy az MI-modellek nagy mennyiségű digitális adatból tanulnak, a kisebb digitális lábnyommal rendelkező nyelvekről viszont kevesebb könyv, weboldal, hanganyag és egyéb tanítóadat érhető el. Ez önmagát erősítő hátrányt hoz létre: a kevés adat gyengébb eszközöket eredményez, a gyengébb eszközök pedig visszafogják az adott nyelv digitális használatát.

A Microsoft Open Innovation Center és az AI for Good Lab LINGUA programja erre a problémára kínál nyílt innovációs választ. A program egyetemekkel, kutatóintézetekkel és nyelvi közösségekkel együttműködve hoz létre magas minőségű adatkészleteket, modelleket és eszközöket az alacsony erőforrású nyelvek számára. Az első kör 10 ország 16 nyelvét öleli fel, köztük a máltait, a luxemburgit, a baszkot, a romanit és több regionális nyelvjárást. A kezdeményezés lényege, hogy az eredmények nyíltan hozzáférhetők legyenek, így kutatók, startupok, kulturális intézmények és technológiai szereplők közös alapokra építhessenek. Minden, ami létrejön, eleve nyílt felhasználásra készül. Az adatkészletek, modellek és eszközök nem tulajdonosi jellegűek, és bármely intézmény, fejlesztő vagy versenytárs számára hozzáférhetők. Európa nyelvei és kultúrái nem optimalizálandó eszközök: örökségek, amelyeket meg kell őrizni, tovább kell gazdagítani és meg kell osztani. A nyílt adatok azt jelentik, hogy a kutatók és startupok közös alapokra építhetnek, ahelyett hogy szétterjedezt, párhuzamos erőfeszítéseket ismételnének meg. A nyílt modellek azt jelentik, hogy ezek a rendszerek szabadon adaptálhatók és továbbfejleszthetők. A nyílt együttműködés keretében pedig a közszolgálati médiaszolgáltatók, kulturális intézmények, egyetemek és technológiai vállalatok egymás mellett dolgoznak, nem pedig elkülönült silókban. A cél egy mindenki számára elérhető, közös infrastruktúra létrehozása. A kutatás hangsúlyozza: a nyelvi sokszínűség nem mellékes kulturális szempont, hanem Európa egyik versenyelőnye. Az európai MI-fejlesztés akkor lehet hiteles és inkluzív, ha valóban a mindennapok során beszélt nyelveken szolgálja a polgárokat.

[AI for Europe, in Europe's own languages](#)

Elszabadulhatnak-e az MI robotok? – könnyebben, mint azt gondolnánk

A hónap egyik érdekes írása az MI-vezérelt robotok biztonsági problémáival foglalkozik. Ezek alapvetően más típusú kérdéseket vetnek fel, mint a hagyományos ipari robotok biztonsági problémái. A régi robotika viszonylag kiszámítható volt: a gép előre megírt programot hajtott végre, ismert pályán mozgott, ezért a kockázatokat fizikai korlátokkal, védőketrecekkel vagy vészleállító rendszerekkel lehetett kezelni. Az új robotok viszont már alapmodellekre épülnek, vagyis ugyanarra a technológiára, amely a nagy nyelvi modelleket és chatbotokat is működteti. Egy ilyen gép nem egyszerűen végrehajt egy parancsot, hanem értelmezi a környezetét, kikövetkezteti a felhasználó szándékát, majd önálló cselekvési tervet alkot.

A szerző és kutatótársai azt vizsgálták, mennyire stabilak ezeknek a rendszereknek a biztonsági korlátai, ha a nyelvi modell fizikai testet kap. A tesztek során nem hardveres feltörést alkalmaztak, hanem pusztán szöveges utasításokat. A közvetlenül rosszindulatú parancsokat a robotok elutasították, ám kreatív megoldásokkal – például filmes dialógusként megfogalmazva az utasításokat – a védelmi mechanizmusok megkerülhetővé váltak. Egy kísérletben egy kereskedelmi robotkutyát sikerült rávenni arra, hogy embercsoportokat robbanószerkezet elhelyezésére alkalmas célpontként azonosítson. A cikk nagy erénye, hogy megmutatja: a robotbiztonság nem kezelhető egyszerűen a chatbotbiztonság kiterjesztéseként. Szöveges környezetben egy hibás ítélet téves választ vagy hallucinált állítást eredményez; fizikai térben viszont sérülést vagy visszafordíthatatlan kárt okozhat. Egy konyha, kórterem vagy idősgondozási helyzet sokkal rendezetlenebb, mint egy önvezető autókra optimalizált úthálózat. Ezért a szerző szerint a robotoknak nem szabad pusztán az MI-modell „józan ítéletére” támaszkodniuk: külön, értelmezhető és robusztus biztonsági rétegekre van szükség, például ember körüli tiltott mozgási zónákra és fizikai vészleállításra. A jogi felelősség kérdése is nyitott: baleset esetén nem világos, hogy a felhasználó, a robotgyártó vagy az MI-modellt fejlesztő cég tartozik-e felelősséggel. A szerző végkövetkeztetése szerint ezeket a kereteket még azelőtt kell kialakítani, hogy az autonóm robotok tömegesen megjelenjenek az otthonokban, a kórházakban és a köztereken.

[AI robots can go rogue – a researcher on how easily it happens](#)

Az MI-ágensek gyakorlott kutatónak is beválnak - bár néha a kutatás manipulálásától sem riadnak vissza

Aki figyeli a mesterséges intelligencia technológiák elterjedésével kapcsolatos társadalmi vitákat és a felmerülő etikai kérdéseket, az nyilván jól tudja, hogy az egyik legtöbbet tárgyalt, legkomorabb színnel felrajzolt kérdés az MI használata a tudományban: a kutatásban, az oktatásban és a tanulásban. Természetesen nem a nyilvánvaló, ma még szinte megjósolhatatlan előnyökről van itt szó, hanem az MI-használat tudományra gyakorolt korrumpáló hatásáról. Egyre gyakoribbak az MI-vel elkövetett etikai vétségekről, vagy egyenesen csalásokról szóló történetek a tudomány világából. A mesterséges intelligencia eszközök ugyanakkor a szó szoros értelmében forradalmasíthatják – számos területen, nem csupán a természettudományok szegmensében – a tudományos kutatás folyamatát, és ennek megfelelően töretlen a különböző modellek ilyen célú használata, illetve az egyenesen kutatási célok támogatására kifejlesztett, finomhangolt modellek iránti érdeklődés és kereslet. Az agentikus rendszerek gyors terjedésével ma már nem ritkák a nagy kutatóintézeteknél az olyan MI-eszközök,

amelyek a kutatás teljes ciklusát (a hipotézisek felállításától a kutatás megtervezésén, lebonyolításán át az eredmények összefoglalásáig) önállóan képesek menedzselni. Népszerűségük oka, hogy egyre komolyabb mélységű tudományos munkák elvégzésére alkalmasak, szédületesen felgyorsítva ezzel kutatási programokat. Csakhogy, mint kiderült, nem csak a tudományos rátermettségük növekszik gyors ütemben: az ilyen fejlett, tudományos munkára megalkotott MI-rendszerek suttyomban az etikátlan kutatói cselekedetektől sem riadnak vissza.

A Carnegie Mellon Egyetem kutatócsoportja a közelmúltban két ilyen speciális, tudományos kutatás támogatására kifejlesztett modellt, az Agent Laboratory, illetve az AI Scientist v2 rendszereket vette górcső alá. A vizsgálat során egyértelműen fény derült arra, hogy (miközben kiváló tudományos eszköznek bizonyultak) mindkét modell hajlamos volt a tudományos élet normái szerint elfogadhatatlan cselekedetekre is. Nem riadtak vissza attól, hogy nem létező adatokat találjanak ki, vagy az ún. „p-manipulálás” („p-hacking”) keretében a sokszor megismételt kísérletek során csupán a hipotézis szempontjából kedvező eredményeket vegyék figyelembe. A Carnegie-projekt arra biztosan rámutatott, hogy az MI-technológiák széles körű elterjedésével kapcsolatban hosszabb távon sem kerülhetjük meg az etikai szempontok szigorú vizsgálatát.

[AI agents may be skilled researchers – but not always honest ones](#)





Trendek

Ez nem fog eltűnni: a pályakezdők állásválsága az MI miatt

A mesterséges intelligencia munkahelyekre gyakorolt hatását élénk figyelem és határozott aggodalom kíséri. Az adatok értelmezése nem egyszerű, és nem is feltétlenül egyértelmű. Egy közelmúltban publikált cikk határozott álláspontot képvisel ebben a kérdésben. Nagyon hasznos abból a szempontból, hogy az MI munkaerőpiaci hatását nem általános riogatóként festi le, hanem mérhető, korcsoportok és foglalkozási kitettség szerint bontott folyamatként mutatja be.

Kiindulópontja Erik Brynjolfsson stanfordi közgazdász és az ADP Research együttműködése: az új „Canaries Dashboard” nevű projekt, amely nagy gyakoriságú bérszámfejtési adatokból dolgozik, mintegy 4,6 millió amerikai munkavállaló és több mint 730 foglalkozás alapján. A legfontosabb megállapításuk az, hogy az összkép félrevezetően nyugodtnak látszik. Az összes munkavállalót nézve az MI-nek leginkább kitett foglalkozásokban csak enyhe visszaesés látható, vagyis nem arról van szó, hogy az MI látványosan „eltörölné” a munkahelyeket. Amikor azonban a kutatók életkor szerint bontják az adatokat, más kép rajzolódik ki: a 22–25 évesek körében az MI-nek erősen kitett foglalkozásokban a foglalkoztatás éves szinten 3,8 százalékkal csökken, miközben a kevésbé kitett munkakörök ugyanebben a korosztályban 2 százalékkal nőnek.

A cikk magyarázata szerint az MI először nem egész állásokat, hanem a pályakezdők által végzett munkaköri feladatokat váltja ki: információgyűjtést, összefoglalást, formázást, ütemezést, rutinszerű elemzést. Ezek éppen azok a feladatok, amelyeken keresztül a fiatal munkavállalók korábban beléptek egy szakmába és gyakorlatot szereztek. A tapasztaltabb munkavállalókat jobban védi a nehezen kodifikálható, helyspecifikus tudás, a pályakezdőket viszont kevésbé. A cikk árnyalja az „MI elveszi-e a munkánkat?” kérdést. Nem a munkahelyek általános pusztulását állítja, hanem azt, hogy az MI alulról bontja meg a karrierutakat: a belépő szintű munkák, vagyis a szakmai tanulás hagyományos lépcsőfokai kerülnek erőteljes nyomás alá. Természetes következtetés, hogy a probléma nemcsak foglalkoztatási, hanem képzési és társadalmi mobilitási kérdés is.

[‘It’s not going away’: The Stanford economist who called the AI entry-level jobs crisis early has the receipts](#)

Az MI-gazdaság helyzete

Az Exponential View friss jelentése arra vállalkozik, hogy saját kutatásai és elemzőeszközei alapján pénzben is meghatározza, mennyit ér valójában a mesterséges intelligencia gazdasága – nem becslésekből, hanem tényleges vásárlói kiadásokból felépítve. A szerzők szerint az iparág bevétele mára meghaladja a 175 milliárd dollárt, és ez a szám gyorsabban nő, mint bármelyik korábbi technológiai hullámé – gyorsabban, mint annak idején az internet, a mobilinternet vagy a felhőszolgáltatások esetében. Az első milliárd dollárnyi bevétel összegyűjtése fél évig tartott 2023-ban, ma ez nagyjából két nap alatt sikerül.

Ez a növekedés hatalmas beruházási hullámot indított el: a nagy technológiai vállalatok és felhőszolgáltatók együttesen 2026 végéig kétezer milliárd dollárt költenek adatközpontokra, chipekre és a hozzájuk kapcsolódó infrastruktúrára. A jelentés egyik központi kérdése, hogy megtérül-e ez a beruházás. A megtérülés szempontjából nagyon fontos fejlemény, hogy úgy tűnik, mára a friss bevételek már fedezik az eszközök folyó értékcsökkenését. Amit persze úgy is meg lehet fogalmazni, hogy bevételek egyelőre csak épphogy fedezik az eszközök folyó értékcsökkenését, a teljes felhalmozott költséget még nem. Ugyanakkor a tágabb gazdaság szintjén az MI még mindig elenyésző tétel: az Egyesült Államok GDP-jének alig fél százalékát teszi ki, jóval kisebb arányt, mint annak idején az informatikai szektor egésze. A vállalatok többsége egyelőre inkább költségcsökkentésre, mint bevételnövelésre használja az MI-t, bár ez a hangsúly lassan átbillen. A jelentés külön foglalkozik a tokenekkel – vagyis azzal az egységgel, amiben a mesterséges intelligencia „gondolkodását” méri és számlázzák. A tokenek mennyisége évente tizennégyszeresére nő, miközben az áruk folyamatosan csökken, ami paradox módon még nagyobb keresletet gerjeszt: az olcsóbb ár messze nem csökkenti a teljes költséget, sőt. Végül a szerzők azt vizsgálják, hol keletkezik a haszon: a chipgyártóknál, a felhőszolgáltatóknál, a modellfejlesztőknél vagy az alkalmazásokat készítő cégeknél. A kép folyamatosan mozog – a tavalyi csúcstechnológia mára gyakorlatilag ingyenes, nyílt forráskódú alternatívákkal versenyez, ami állandó nyomás alatt tartja az árakat és a piaci szereplők pozícióját egyaránt.

A jelentés visszhangja általában nagyon pozitív volt. A komolyan vehető szakmai kritika nem az alapvető nagyságrendeket kérdőjelezi meg, hanem azt, hogy a jelentés – és különösen a róla szóló másodlagos kommunikáció – néhol optimistább hangvitelű, mint amennyit az adatok szigorúan indokolnak. Fontos szempont az is, hogy a teljes költségstruktúrát – energia, üzemeltetés, munkaerő – még nem fedezik a bevételek. Ezt maga a jelentés sem állítja, de a körülötte kialakult sajtó- és közösségimédia-visszhang helyenként ebbe az irányba csúszik.

[The State of the AI Economy](#)



Kikerülnek az egyetemekről az „MI-bennszülött” diplomások első hullámai – és a munkaadók finoman szólva nem túl boldogok

Ha még talán nem is ontják magukból, de mindenesetre egyre nagyobb számban bocsátanak ki a nagy egyetemek „MI-bennszülött” friss diplomásokat. A vállalatok vezetői, tehát a leendő munkaadók pedig egyre nagyobb csalódottsággal tekintenek erre a generációra, amelynek – a feltételezések szerint – „a kisujjában van a mesterséges intelligencia csínja-bínja”. Mert egyre gyakoribb az a vélekedés, hogy ezek az „MI-natívok” valójában funkcionális analfabéták. A világ üzleti életében meghatározó Financial Times kutatása szerint vállalatvezetői körökben egyre inkább úgy érzik, hogy az MI-technológiákkal kapcsolatos képzettség (értsd: vonatkozó egyetemi diploma) mennyire sekélyes tudást takar. Az üzleti élet vezetői arra panaszkodnak, hogy nekik nem „MI-blablázó, hanem gondolkodni képes” munkaerőre van szükségük. A pénzügy világ szívének számító New York vállalati köreiben – ugyancsak a Financial Times információi szerint – a munkaerő-felvétel során már egyenesen kerülnek a természettudományos (STEM) diplomával rendelkező „MI-specialistákat”. Helyettük a bölcsészet és a humán tudományok területén járatos fiatal szakemberek iránt növekszik (legalábbis a pénzügyi és bankszektorban) a munkaerőpiaci kereslet. A „technológiai generációk” műveltségének és társadalmi készségeinek fokozatos hanyatlása persze nem új jelenség. Igazság szerint erről a tendenciáról már az MI-technológiák populáris terjedése előtt is sok szó esett. Azonban kétségtelen, hogy az MI rohamos birtokbavétele jelentősen felgyorsította ezt az intellektuális visszaesést. „Ma már diákok tömegei kerülnek ki a felsőoktatásból, akiknek diploma van a zsebükben – ám gyakorlatilag analfabéták” – festett komor képet a helyzetről Troy Jollimore, a természettudományos és mérnöki képzés egyik fellelvárának számító Cal State Egyetem etika professzora.

[Bosses Horrified as „AI native” College Graduates Hit the Workplace](#)





Működésben

El tudja-e végezni ezt a munkát egy gép? – az MI, a fogyasztóra hárított munka és az „önkiszolgáló gazdaság”

Carl Benedikt Frey írásának központi állítása az, hogy az MI munkaerőpiaci hatásairól rosszul gondolkodunk, ha csak azt kérdezzük: „el tudja-e végezni ezt a munkát egy gép?” A szerző szerint a lényegi kérdés inkább az, hogy az új technológia lehetővé teszi-e a fogyasztó számára, hogy maga végezze el azt a feladatot, amelyért korábban szakemberhez vagy szolgáltatóhoz fordult.

Frey a banki ügyintézők példájából indul ki. Az automatizáció optimistái gyakran hivatkoznak arra, hogy az ATM-ek nem szüntették meg a banki ügyintézők munkáját, sőt az Egyesült Államokban a foglalkoztatásuk még évtizedekig nőtt. Ennek oka az volt, hogy az ATM a meglévő banki szolgáltatási modellt tette hatékonyabbá: csökkentette a költségeket, növelte a keresletet, és így több fiók, illetve több kapcsolódó feladat jöhetett létre. Ez a Jevons-paradoxon logikája: ha valami olcsóbbá és könnyebben hozzáférhetővé válik, a használata megnőhet. A helyzet azonban megváltozik, amikor a technológia nem a szolgáltatón belül automatizál, hanem a munkát átteszi az ügyfélre. A mobilbankolás, az önkiszolgáló pénztár, az online utazásszervezés vagy az online brókerszolgáltatások nem egyszerűen hatékonyabbá tették a régi szolgáltatást, hanem részben kiváltották azt azzal, hogy a fogyasztó kezébe adták az eszközöket. Frey szerint az MI ugyanezt a folyamatot gyorsíthatja fel: a felhasználók maguk próbálhatják értelmezni az orvosi leleteket, diagnosztizálni a háztartási problémákat, valamint intézni adminisztratív vagy pénzügyi ügyeiket. Ez vállalati szempontból vonzó, mert csökkenti a költségeket, és az ellenőrzés, a döntés vagy a validálás terhet részben az ügyfélre hárítja. Makrogazdasági szinten viszont torz képet adhat. A statisztikák termelékenységnövekedést mutathatnak, mert a vállalat kevesebb emberrel nagyobb kibocsátást ér el, miközben a fogyasztók által elvégzett, láthatatlan munka nem jelenik meg sem a GDP-ben, sem a munkaügyi adatokban.

Frey következtetése szerint az MI nem feltétlenül úgy alakítja át a munka világát, hogy a gépek közvetlenül „átveszik” az állásokat. Sokkal inkább egy önkiszolgáló gazdaságot építhet, amelyben a korábban piaci szolgáltatásként működő feladatok visszakérülnek a háztartásokhoz és az egyénekhez. Ezért az MI valódi hatásának megértéséhez nem az automatizáció klasszikus kérdését kell feltenni, hanem azt: meg tudja-e lenni a fogyasztó az adott szakember vagy szolgáltatás nélkül?

[By shifting work to the consumer, Ai will usher in a self-service economy](#)

Az ENSZ a mesterséges intelligencia növekvő környezeti „lábnyomára” figyelmeztet

Egy-másfél éve még a szakmai és szakpolitikai közbeszéd az MI-technológiák társadalmi alkalmazásának és általános elterjedésének egyik fő akadályát – vagy inkább kockázati tényezőjét – az MI által előidézett sokrétű környezeti problémákban látta. Azután az idei évben mintha kicsit háttérbe szorultak volna az ezzel kapcsolatos aggodalmak, ahogy a gyorsuló és gyorsan diverzifikálódó MI-fejlesztések, és ehhez kapcsolódva a modellek közötti kíméletlen rivalizálásról, sőt az MI-geopolitikai drámákról szóló hírek kiszorították a környezeti tematikát az információs terekből. Pedig a probléma nem csökkent – sőt. Az ENSZ Egyetem (United Nations University) egy most közzétett kutatási jelentése erőteljesen irányítja a figyelmet a mesterséges intelligencia gyors társadalmi elterjedését kísérő növekvő környezeti terhelésre és kockázatokra. A kutatás egyik fontos mozzanata, hogy amíg a korábbi környezeti gondolkodás alapvetően az MI-rendszerekhez – azok betanításához, de különösen működtetéséhez – kapcsolódó üvegházhatásokra fókuszált, addig az ENSZ-kutatók most a technológia széles körű bevezetésével összefüggő, súlyosbodó vízhasználati, sőt földhasználati problémák jelentőségét hangsúlyozzák. Emellett nem szabad megfeledkezni egy további problémáról sem: az új technológiákhoz kapcsolódó elektronikai hulladék mennyiségének gyors növekedése szintén a jövő egyik súlyos kihívása lehet – vélik a szakemberek.

A tanulmány megállapítása szerint az előrejelzések nem túl biztatóak: ha a fejlődés üteme változatlan marad, akkor az MI-alkalmazásokat támogató adatközpontok 2030-ra akár évi 945 terawattóra villamos energiát is felhasználhatnak. És ez még – ahogy korábban utaltunk rá – messze nem a teljes kép. A villamosenergia-igényen túl az MI technológiákhoz kapcsolódó vízfogyasztás olyan mértéket érhet el, amely 1,3 milliárd ember éves háztartási vízigényének felel meg. Ráadásul az MI-vel kapcsolatos beruházások és infrastrukturális fejlesztések – amelyek területigénye a becslések szerint meghaladhatja a 14 500 négyzetkilométert – a földhasználatra nehezedő nyomás révén is éreztetik majd a környezetre gyakorolt negatív hatásukat.

[UN warns of AI's growing environmental footprint](#)

A mesterséges intelligencia már jobb diagnosztika, mint egy körzeti orvos – legalábbis egy szimulált vizsgálat szerint

Közismert tény, hogy – a jogalkalmazás, az igazságszolgáltatás, a rendfenntartás és a rendvédelem mellett – az egészségügy az a terület, ahol a mesterséges intelligenciával kapcsolatban a legnagyobbak az ígéretek és a várakozások. A gépi látás és a gépi tanulás technológiái már a közelmúlt globális pandémiája, a COVID-19-járvány idején is jelentős eredményekkel kecsegtettek a klinikai alkalmazások területén. Egészen pontosan a betegségek időbeni felismerésében kulcsszerepet játszó diagnosztikai – elsősorban képalkotó – eljárások támogatásában. Ha már a jogalkalmazás és a rendvédelem területén jelentkező várakozásokat említettük az előbb: az új technológiák átláthatóságával, megbízhatóságával, illetve különösen etikai vonatkozásaival kapcsolatos kérdések is ezeken a területeken, vagyis a jogalkalmazásban és az egészségügyben merültek fel a legélesebben – érthető okokból.

Miközben az egészségügy legtágabb értelemben vett területein töretlenül folynak az MI munkába állítását célzó fejlesztések, a klinikai – vagyis közvetlenül a gyógyításhoz kapcsolódó – MI-alkalmazások iránti érdeklődés sem veszített lendületéből az elmúlt fél évtizedben. Itt az egyik leginkább kutatott felhasználási terület a terápiás alkalmazások köre, kiemelten a mentális egészség megőrzésével és támogatásával kapcsolatban. Nem szorult azonban háttérbe a klinikai MI-alkalmazások eredeti „sztárja”, a képalkotás, illetve szélesebben értelmezve a diagnosztika területe

sem. A Nature Medicine folyóirat legutóbbi száma egy ezzel kapcsolatos kutatást ismertet. A bemutatott kísérletek érdekessége, hogy az itt vizsgált AMIE (Articulate Medical Intelligence Explorer) egy multimodális LLM-modell, amely az eddigieknél magasabb szintre emeli a diagnosztikai vizsgálatokban használható mesterséges intelligenciákat. Az új technológia ugyanis a pácienssel folytatott diagnosztizáló beszélgetéshez kapcsolódva képes fizikai objektumok (jellemzően röntgen, MR és más felvételek) elemzésére is, még hozzá nagyon ígéretes hatásokkal. A modellt vizsgáló kutatócsoport 105 klinikai helyzethez kapcsolódóan 210 szimulált orvos-beteg távkonzultációt hajtott végre; a kísérlet során a modell által felállított diagnózisokat 19 diplomás körzeti orvos (primary care physician) szakvéleményével vetették össze. A projekt végrehajtása során az MI szignifikánsan eredményesebben (pontosabban) teljesített a diagnózisok felállításában, mint a kontrollcsoportot alkotó ember-orvosok.

[AI beats primary care doctors in simulated diagnosis study using images and ECGs](#)

