



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

2025 november

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

- A munka a jövőben az emberek, MI-ügynökök és robotok alkotta partnerségen fog alapulni
- Az MI túlságosan nagy kockázatot jelent ahhoz, hogy biztosítást lehessen rá kötni
- Hogyan vált az általános mesterséges intelligencia elképzelése korunk legjelentősebb összeesküvés elméletévé?
- Az ausztrál rendőrség erre a célra fejlesztett nagy nyelvi modellel fejtené meg az online zaklatók szleng-beszédét

Trendek

- Az ember, mint autonóm algoritmus-nyájak pásztora – Technológiai trendek 2025-2026
- Az OpenAI cég az egészségügyi (fogyasztói) alkalmazások piaca felé kacsintgat
- A mesterséges intelligencia „társak” dinamikusan növekvő új iparága

Működésben

- Az MI igényei miatt félévente kell megduplázni a Google számítási kapacitásait
- A résztvevőket helyettesítő chatbotok miatt változóban az online konferenciák és megbeszélések műfaja
- LLM-buborékban élünk – nem pedig MI-burokban
- Az energiára szomjazó MI-adatközpontok: helyzetkép az USA-ból





Etika és jog

A munka a jövőben az emberek, MI-ügynökök és robotok alkotta partnerségen fog alapulni

A mesterséges intelligencia átalakítja a munka világát. Előnyeinek kiaknázása új készségeket és annak újragondolását igényli, hogyan működjenek együtt az emberek az intelligens gépekkel. A McKinsey cég kutatási jelentése az „ügynökök” és a „robotok” kifejezéseket tág, gyakorlati értelemben használja: előbbin minden olyan gépet értenek, amely képes nem fizikai jellegű munkafolyamatok automatizálására, utóbbin pedig azokat, amelyek fizikai tevékenységeket automatizálnak. Számos különböző technológia látja el ezeket a funkciókat – némelyik MI-alapú, mások nem –, és a köztük húzóó határvonalak folyamatosan változnak és elmosódnak. A jövő munkájának alapja az emberek, az ügynökök és a robotok közötti partnerség lesz, mindezt át fogja hatni a mesterséges intelligencia. A mai technológiák elméletben a jelenlegi amerikai munkakörökben végzett munkaórák több mint felét automatizálhatnák. Ez a kijelentés a munka átalakulásának mélységére utal, de nem a munkahelyek sorsának (megszűnésének) előrejelzése. Az elterjedés fokozatos lesz, a folyamat során egyes munkakörök zsugorodni fognak, mások bővülni vagy átalakulni, miközben újak jönnek létre – és a munka egyre inkább az emberek és az intelligens gépek együttműködésére épül majd. Az emberi készségek túlnyomó része megmarad, de másképp alkalmazzuk majd őket. A munkáltatók által ma keresett készségek több mint 70 százaléka jelen van az automatizálható és a nem automatizálható feladatokban egyaránt. Ez az átfedés azt jelenti, hogy a legtöbb készség továbbra is releváns marad, de az alkalmazásuk módja és közege változik. A kutatók által kidolgozott Készségváltozási Index (Skill Change Index) rávilágít arra, hogy mely készségek lesznek leginkább és legkevésbé kitéve az automatizáció hatásainak a következő öt évben. A digitális és információfeldolgozási területeken várhatók a legnagyobb változások, míg a gondoskodással összefüggő készségek esetében a hatás mérsékeltebb. Az MI-eszközök alkalmazásának és kezelésének képessége iránti kereslet két év alatt hétszeresére nőtt, gyorsabban, mint bármely más készség iránti igény az amerikai álláshirdetésekből. A felfutás iparágakon átívelően megfigyelhető, és nagy valószínűséggel egy még jelentősebb átalakulás kezdetét jelzi. Összességében 2030-ig mintegy 2,9 billió dollárnyi új értéket teremthet az, ha a szervezetek felkészítik munkavállalóikat az új technológiákra. Ehhez újra kell tervezni a munkafolyamatok egészét – nem csupán az egyes feladatokat – az emberek, ügynökök és robotok integrált együttműködése alapján.

[Agents, robots and us: Skill partnership in the age of AI](#)

Az MI túlságosan nagy kockázatot jelent ahhoz, hogy biztosítást lehessen rá kötni

A nagy biztosítótársaságok nagyon óvatosak az MI-rendszerek működéséből származó kockázatokkal kapcsolatos hozzáállásukban. A vállalatok számára nyújtott biztosítások esetében igyekeznek ezeket a kockázatokat kizárni, azaz a vállalatok által készített MI-rendszerek esetleges hibás működése következtében fellépő kárigényekre nem nyújtanak fedezetet. A vállalatok viszont adott esetben több milliárd dolláros igényekkel nézhetnek szembe az egyre gyorsabban fejlődő technológia miatt. A nagy biztosítók egyre jobban tartanak a mesterséges intelligenciához köthető, potenciálisan rendszerszintű káreseményektől, ezért sorra kérik az engedélyt az MI-kockázatok (amelyek például a chatbotok, ügynökök alkalmazásából származnak) kizárására a vállalati biztosítási kötvényekből. A generatív modellek kiszámíthatatlansága és a felelősségi viszonyok tisztázatlansága miatt a fedezet drámaian szűkülhet, miközben a technológiát alkalmazó cégek száma gyorsan nő. Egy európai szakértő szerint a biztosítók egyre inkább úgy látják, hogy az MI-modellek outputjai túl kiszámíthatatlanok és átláthatatlanok ahhoz, hogy biztosíthatók legyenek. Túlságosan fekete a doboz – mondhatnánk. Több nagy visszhangot kapott incidens – például MI által okozott rágalmazás, hamis chatbot-ajánlat vagy mélyhamisításon alapuló pénzügyi csalások – már most is jelentős pénzügyi károkat okozott. Jelenleg az is nehézséget jelent, hogy megállapítsák, ki a felelős. Ezzel szemben az MI-kockázatok potenciálisan sok szereplőt érinthetnek – fejlesztőket, modellezőket, végfelhasználókat. Emiatt az MI-kockázatok piaci hatása exponenciális is lehet. Az iparág elbírna még egy 400–500 millió dolláros veszteséget is, amelyet egy ügynökszerű MI-rendszer által okozott hibás árazás vagy orvosi diagnózis okoz egyetlen vállalatnál. Amit nem tudnának elviselni, az, ha egy MI-szolgáltató hibája ezer vagy tízezer esetben okoz veszteséget, mert ez már egy rendszerszintű, korrelált, aggregált kockázat. Erre a helyzetre egyre több biztosító kiegészítésekkel, szűkítésekkel vagy részleges fedezetekkel válaszol, így próbálják csökkenteni a bizonytalanságot. A biztosítási brókerek és az ilyen esetekre szakosodott ügyvédek attól tartanak, hogy a biztosítók a bíróságon kezdik majd vitatni az igényeket, amikor az MI-vezérelt károk jelentősen megszaporodnak.

[AI is too risky to insure, say people whose job is to insuring risk](#)

Hogyan vált az általános mesterséges intelligencia elképzelése korunk legjelentősebb összeesküvés elméletévé?

Az elképzelés, miszerint a gépek értelmi képességei el fogják érni az emberi szintet – vagy akár okosabbak is lehetnek – egy egész iparágat rabul ejtett. De ha jobban megnézzük, látni fogjuk, hogy ez egy mítosz, amelynek elterjedése és fennmaradása sok tekintetben megegyezik az összeesküvés-elméletek esetében tapasztalt jelenségekkel. A mesterséges általános intelligenciában (artificial general intelligence, AGI) való hit napjaink egyik meghatározó tényezője az MI-fejlesztők élcsapatának és az őket finanszírozó kockázati tőkebefektetők (a szilícium-völgyi elit) világképének. Az AGI körül tapasztalható hatalmas felhajtás, az ebből fakadó apokaliptikus félelmek és a megváltáshoz hasonló ígéretek sok szempontból hasonlítanak az összeesküvés-elméletekhez. Az AGI hívei rendkívüli jövőképet festenek: egyrészt az AGI eljövetele meg fogja oldani a társadalmi problémákat, felszámolja a szegénységet, leküzd a betegségeket, másrészt az apokaliptikus jövőképük szerint el fogja pusztítani az emberiséget. Ezt a nagyon erős narratívát misztikus várakozások, a világ nemsokára bekövetkező gyökeres átalakulásának ígérete, egyben a kiforrott tudományos bizonyítékok hiánya jellemzi. Végső soron az AGI egy modern hitrendszerként működik, amely a technológiát kínálja egyetlen megváltásként a világ problémáira. Az AGI mint közeli, elkerülhetetlen jövő elképzelése eltúlzott,

megalapozatlan. A definíciója rendkívül homályos: nincs konszenzus arról, mit jelent pontosan AGI-t létrehozni, mik ennek a mérhető kritériumai, különösen az emberi képességek szintjét elérő, sokoldalú, adaptív intelligenciával rendelkező gépek esetében. Ez a bizonytalanság persze lehetővé teszi, hogy a célokat (dátumok, képességek, ígéretek) rugalmasan lehessen változtatni. Nagyon nagy problémát okoz(hat) az, hogy a világunkat elpusztító AGI-val való terméketlen hadakozás eltereli a figyelmet a most létező és rohamosan fejlődő MI-rendszerek számos problémájának szabályozásáról. A felhajtás segítheti a technológiai és üzleti befektetéseket, de nem feltétlenül jár együtt tudományos megalapozottsággal – így az eshetőség hangsúlyozása (egyszer be fog következni) üzleti eszközzé válik. A rendkívüli képességek, rendkívüli lehetőségek, rendkívüli veszélyek súlykolása egyben igazolja a rendkívüli befektetéseket és a fejlődés útjában álló szabályozás eltávolítását. Mindez a társadalmi befolyás növelését is szolgálja. Érdekes a jelenség részletes történeti áttekintése is. Hogyan és mikor merült fel először az emberi képességeket felülmúló mesterséges intelligencia fogalma, kik voltak a megfogalmazói? A MIT Technology Review cikke hangsúlyosan tárgyalja ezeket a kérdéseket az informatikai ipar fejlődésének utóbbi kb. 15 évének történéseibe ágyazva. Felvázolja, hogy hogyan vált szélsőséges elképzelésből uralkodó tantétellé az AGI, továbbá részletesen bemutatja a Szilícium-völgy meghatározó vállalkozóinak, befektetőinek szerepét ebben a folyamatban.

[How AGI became the most consequential conspiracy theory of our time](#)

Az ausztrál rendőrség erre a célra fejlesztett nagy nyelvi modellel fejtené meg az online zaklatók szleng-beszédét

A mesterséges intelligencia technológiák egyik „kísérleti” alkalmazási területe, hagyományosan, a belbiztonság (rendőrségi munka) és az igazságszolgáltatás, annak ellenére, hogy az MI (kiforratlansága, precedens nélkülsége, fekete doboz jellege folytán) egész sor etikai, személyiségi jogi, jogalkalmazási, sőt jogfilozófiai problémát vet fel. A magyarázat egyszerű: a közbiztonságon őrködő rendőri erők leterheltsége óriási, akárcsak a jogszolgáltatást végző bíróságoké. Az ausztrál rendőrség most ezen a sikamlós terepen, ám teljesen újszerű felhasználásra épít egy MI-eszközt. A Microsofttal összefogva fejlesztett technológia nem a szokványos arcfelismerő, viselkedés-előrejelző, esetleg döntés-támogató szerepkörökben igyekszik bekapcsolódni a rendőri munkába. A speciális LLM-et arra készítik fel, hogy az online bűnözők sajátos jelnyelvét (belső szlengkifejezéseit, vagy éppen saját használatú jeleit, emotikonjait) segítsen megfejteni. Általános felismerés, hogy a közösségi média platformok egyes digitális bűnözői magatartások, tevékenységek melegágyaivá váltak. Világszerte terjed a kiskorúak sérelmére elkövetett bűncselekmények összetett csoportja (amelyet a nyugati világban a kissé semmitmondó „zaklatás” fogalommal szoktak jelezni), és immár dokumentálható új jelenséggént kell számolni az online terekben működő, ún. „bűnöző-influenszerek” aktivitásával is. Szakemberek azt remélik az új fejlesztéstől, hogy a rendőrség időben képes lesz felfigyelni a – kiskorúak biztonságát fenyegető – veszély jelekre, még ha azok „virág nyelven” is lennének megfogalmazva.

[Australian police create AI tool to decode predators' slang](#)





Trendek

Az ember, mint autonóm algoritmus-nyájak pásztora – Technológiai trendek 2025-2026

Az Európai Adatvédelmi Biztos (EDPS, European Data Protection Supervisor) az uniós intézmények és szervek adatvédelmi felügyeletéért felelős hatóság vezetője. Hivatalának éves technológiai előrejelzéseit TechSonar címen jelentetik meg. A TechSonar idei jelentése hat trendet azonosít: az ügynöki (agentikus) MI-t, az MI-alapú digitális társakat, az automatizált vizsgafelügyeletet, az MI-vezérelt személyre szabott tanulást, a kódolási asszisztenseket és a bizalmas adattitkosítást biztosító számítástechnikai megoldásokat (confidential computing). Az agentikus MI olyan mesterségesintelligencia-rendszereket jelöl, amelyek képesek önállóan döntéseket hozni, cselekvéseket végrehajtani és célokat elérni anélkül, hogy folyamatos emberi beavatkozásra lenne szükség. Az MI-alapú digitális társak személyre szabott élmények révén lépnek kapcsolatba az emberekkel és támogatják őket. Az automatizált vizsgafelügyelet az online vizsgák megfigyelésével igyekszik kiszűrni a csalást. Az MI-vezérelt személyre szabott tanulás az egyes tanulók igényeihez igazítja a tananyagot és a tanulási élményt. A kódolási asszisztensek a fejlesztők munkáját segítik a kód megírásában és hibakeresésében. A confidential computing pedig olyan védelmet biztosít, amely a feldolgozás alatt is óvja az adatokat, biztonságos, elkülönített környezetekben végrehajtott számításokkal. Ahogy a mesterségesintelligencia-rendszerek egyre önállóbbá válnak, és egyre mélyebben beépülnek az emberi környezetbe, úgy kerülnek az emberek mindinkább olyan szerepbe, amelyet „MI-ügynökök pásztorának” nevezhetünk: vagyis hátrálépnek a feladatok közvetlen elvégzésétől, és inkább a rendszerek felügyeletére, azok működésének irányítására és hatásainak alakítására összpontosítanak, biztosítva, hogy a rendszerek az emberi értékekhez igazodjanak. A kihívás kétirányú lesz az adatvédelmi, jogvédő szakemberek számára. Először is gondoskodni kell arról, hogy a technológiai fejlődés továbbra is összhangban maradjon az alapvető jogokkal. Másodszor pedig garantálni kell, hogy az MI növekvő autonómiája ne járjon az emberi cselekvőképesség csökkenésével – vagyis azzal, hogy az emberek képesek legyenek önálló döntéseket hozni, irányítani saját cselekedeteiket, és felelősek maradjanak döntéseikért, ahelyett, hogy az algoritmusok határoznák meg vagy korlátoznák őket. Ennek eléréséhez célzott kutatásra és technológiai innovációra lesz szükség, valamint a humán értékek melletti következetes kiállásra, átgondolt irányításra és a különböző szakterületek együttműködésére, elősegítve olyan párbeszédet, amely a digitális átalakulás középpontjába a magánszférát, az elszámoltathatóságot és az emberi méltóságot helyezi. A jelentés tartalmaz egy módszertani ismertetőt is az előrejelzési tevékenységhez.

[TechSonar 2025-2026](#)

Az OpenAI cég az egészségügyi (fogyasztói) alkalmazások piaca felé kacsintgat

Régi felismerés, hogy az emberek számára az egészségük az egyik legfontosabb érték. Pontosabban az egészségükért való aggodás, hiszen az egészség megőrzéséért tett aktív erőfeszítések általában jóval hátrébb kullognak a hétköznapi ember napi prioritásaiban. Mindenesetre az egészségügy széles területé az egyik legfontosabb piac, és nem csupán a vállalati szegmensben mozgó szereplők számára. Ez az egyik legjelentősebb fogyasztói piac is: egyrészt mérete, másrészt a hozzá kapcsolódó, előbb említett erős emocionális kötődések, vágyak következtében. A technológiai vállalati kör, és különösen a technológiai óriáscégek – név szerint a Google, az Amazon és a Microsoft – évtizedek óta igyekeznek kiaknázni a fogyasztói egészségügyi adatok értékesítésében rejlő üzleti lehetőségeket. Piaci elemzők úgy vélik, hogy a csillagként emelkedő generatív mesterséges intelligencia vállalatok új lendületet adhatnak ezeknek a törekvéseknek. A nagyok közül többen (itt elsősorban az egyébként is az élen haladó OpenAI-t említik a szerzők) már aktívan kutatják annak a lehetőségeit, hogy az LLM-technológiákra építve hogyan tudnának fogyasztói egészségügyi eszközöket, alkalmazásokat létrehozni. A gondolkodási-fejlesztési irányok között elsősorban a „személyi egészségügyi asszisztenseket”, illetve a „személyes egészségügyi adat-aggregátorokat” szokták említeni. A fogyasztók hosszú ideig jellemző módon a Google keresőjében igyekeztek utánanézni az egészségügyi problémáikhoz, vagy a félelmeikhez kapcsolódó témáknak, információknak. Azonban jól látható az általános tendencia, ahogy a „Google-keresések” gyakorlatából egyre többen pártolnak át a nagy nyelvi modellekkel folytatott beszélgetés-alapú információ felkutatás világába. Ezért okkal tételezhető fel, hogy ez az érdeklődések speciálisabb, szűkebb területeire, így az egészségüggyel kapcsolatos kérdések körére is érvényes lesz majd. A fogyasztói egészségügyi piac befektetői nagy elvárásokkal tekintenek az OpenAI mintegy 800 milliós heti aktív felhasználói bázisára. Sam Altman cége mindenesetre megkezdte a fogyasztói egészségipar néhány kulcsszemélyiségének leszerződését.

[OpenAI is weighing a move into consumer health apps](#)

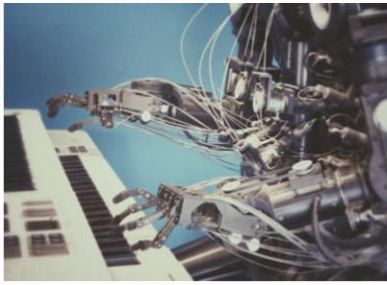
A mesterséges intelligencia „társak” dinamikus növekvő új iparága

Nem sokkal ezelőtt még abszurdnak, de legalábbis a tudományos fantasztikus irodalom lehetséges fordulatának tűnhetett az a gondolat, hogy emberi lények érzelmi kapcsolatot építsenek ki mesterséges intelligencia szoftverekkel. Ma már azonban ez a valóság része – még hozzá dinamikus bővülő, terjeszkedő része. A „kapcsolat”, „kötődés” kifejezést sokféle konkrét tartalommal használja a szakirodalom és a közbeszéd egyaránt. Vannak, akik virtuális „barátnak”, „cimbórának” tekintik az MI-algoritmust (illetve az azon működő chatbotot), mások egyfajta „mentornak” nevezik a chatbotjukat. Egyre többen vannak olyanok, akik „terapeutaként”, azaz pszichológusként, sőt pszichiáterként tekintenek ezekre a modellekre. Sőt, olyanok is akadnak szép számmal, akik „szerelmi kapcsolatot” ápolnak chatbotjukkal. Ez a tömeges jelenség tömegpiaci lehetőségekkel is kecsegtet, és ebben a feltámadó érdeklődésben fejlesztő cégek egész sora látta meg az üzleti lehetőséget. Az elmúlt években több száz speciális „ember-gép érzelmi kapcsolatra” profilírozott applikáció jelent meg, és ezeknek a használata folyamatosan növekszik. A legnagyobb ilyen „társ chatbot” alkalmazás, a Character.ai ma már átlagosan 20 millió körüli aktív havi felhasználót könyvelhet el, de a mögötte haladó Polybuzz forgalma is eléri havonta a 13 millió felhasználót. A jelenség elsősorban a nyugati világot érinti, de érdemes felfigyelni rá, hogy a legnagyobb ilyen „társ chatbot” Kínában, a Maoxiang applikáció is 1,2

millió havi aktív felhasználót számolhat ma már. Ugyanakkor egyértelműen az Egyesült Államokban a legelterjedtebb ez az ember-gép kapcsolatforma: a középiskolás korú amerikaiak 42%-a nyilatkozott úgy, hogy ők maguk, vagy közvetlen ismeretségük már épített ki ilyen „baráti” kapcsolatot chatbottal, ráadásul ezek közel 20%-a a „szerelmi” kategóriába sorolta az algoritmikus viszonyt. A piaci tendenciákhoz igazodva, nem csupán mennyiségben, de minőségben is növekszik a kínálat: az erre a területre szakosodott fejlesztők egyre „emberibb” viszonyt ígérő chatbotokkal lépnek elő.

[A new industry of AI companions is emerging](#)





Működésben

Az MI igényei miatt félévente kell megduplázni a Google számítási kapacitását

A mesterséges intelligencia körüli esetleges iparági buborék (és annak feltételezhető kidurranása) témája egyre gyakrabban jelenik meg a közbeszédben, ugyanakkor a nagy technológiai vállalatok – köztük a Google és az OpenAI – infrastrukturális oldalról épp az ellenkező kihívással szembesülnek: nem tudnak elég gyorsan kapacitást építeni a növekvő MI-igények kiszolgálására. A Google belső tájékoztatása szerint a vállalatnak félévente meg kell dupláznia kiszolgálási kapacitását, és 4-5 éven belül ezerszeresére kell növelnie rendszerei számítási teljesítményét. A vállalat vezetői ugyanakkor hangsúlyozzák, hogy mindezt költség- és energiahatékonysági korlátok között kell megvalósítani, ami komoly műszaki és szervezeti feladat. Az ismertetett kereslet összetételét nem határozták meg közelebbről. Részben a felhasználók valódi igényei állhatnak a háttérben, részben pedig a vállalat szolgáltatásaiba beépített MI-funkciók növelik a terhelést. A Google nincs egyedül, valóságos adatközpont építési láz uralkodik az iparág szereplői között. Az OpenAI például több százmilliárd dolláros beruházással hat új amerikai adatközpontot tervez, hogy kiszolgálja a ChatGPT gyorsan növekvő forgalmát. A kapacitásépítés egyik fő akadály a Nvidia chipgyártási korlátja, amely lassítja az új MI-funkciók széles körű bevezetését is. A cikk nem említi ugyan, de a másik korlátozó tényező a biztonságos energiaellátás megvalósításának nehézségei. A Google három stratégiai irányra épít az adatközponti kapacitás bővítése során: új fizikai infrastruktúra létrehozása, hatékonyabb MI-modellek fejlesztése és saját tervezésű hardverek, köztük a legújabb TPU-generáció alkalmazása. A TPU azaz Tensor Processing Unit, egy kifejezetten mesterségesintelligencia-számításokra tervezett, speciális Google-chip. A vállalat vezetése úgy látja, hogy a túl kevés beruházás nagyobb kockázatot hordoz, mint az esetleges túlkapacitás, ugyanakkor elismerik, hogy az MI-iparág jövője körüli bizonytalanság továbbra is jelen van.

[Google tells employees it must double capacity every 6 months to meet AI demand](#)

A résztvevőket helyettesítő chatbotok miatt változóban az online konferenciák és megbeszélések műfaja

A nagy nyelvi modellek gyors beépülése mindennapjainkba, emberi és munkafolyamatainkba, szinte hetente kínál valami meglepő újdonságot. Az elmúlt hónapban például arról igyekeztünk beszámolni, hogyan forgatja fel fenekestül az LLM-ek rohamosan terjedő használata a munkaerő-piaci személyzeti munkát. Ahol egészen újszerű helyzetet generál az a tény, hogy egyfelől az állásra jelentkezők, másfelől az álláskereső felvételét elbírálók egyre gyakrabban használnak maguk helyett generatív MI-technológiákat. Most egy másik, hasonlóan meglepő térhódítást mutatunk be: azt a jelenséget, ahogyan a mesterséges intelligencia chatbotok fokozatosan átveszik az (emberi) résztvevők helyét a COVID-időszak óta különösen népszerű online megbeszélések, videokonferenciák, illetve más digitális események egész sorában. Szaporodnak a konkrét híradások olyan esetekről, ahol ilyen online megbeszélésre, beszélgetésre a résztvevők egy része egy „személyi asszisztenst” küldött el maga

helyett, vagyis egy MI-chatbotot. Bár a hír első hallásra tényleg hihetetlennek tűnik, az ilyen akciók szervezői azzal érvelnek, hogy a rohanó világban, amikor egy-egy potenciális résztvevő nem tud személyesen időt szakítani egy ilyen online eseményre, de ugyanakkor érdekelt lenne az ott elhangzottakban, a maga helyett „küldött” chatbot-asszisztens jegyzeteket, összefoglalókat készíthet számára az elhangzottakról. Tehát egyfajta „termelékenység növelő” technológiaként segíti emberi alkalmazóját, és éppen a „productivity” növelése az új generatív MI-technológiák egyik nagyon fontos „munkaköre”. A fenti gyakorlattal szemben érvelők azonban azt hangsúlyozzák: a személytelen algoritmusok „részvétele” az alapvetően emberek közötti interakciós területnek tekinthető megbeszéléseken, konferenciákon, aláássa a párbeszédhez nélkülözhetetlen bizalmat.

[Live events rethink rules as AI assistant bots take seat](#)

LLM-buborékban élünk – nem pedig MI-burokban

A HuggingFace vállalat az mesterséges intelligencia színtér egyik meghatározó, sajátos szereplője. Clem Delangue, az MI-platformként és közösségként működő cég népszerű vezetője szerint tévedés azt hinni, hogy MI-buborékban élünk, valójában a nagy nyelvi modellek (LLM) okozta buborékról van szó. A buborék léte vagy nem léte tényleg nagyon fontos kérdés, de nem hiszi, hogy a buborék esetleges kidurranása az egész mesterséges intelligencia jövőjét veszélybe sodorná. A nagy nyelvi modellek (LLM-ek, amelyek a ChatGPT-t, a Geminét és más csevegőrobotokat működtetnek) kapnak aránytalanul nagy figyelmet, és ez a figyelem nem biztos, hogy tartós lesz. Ennek az lesz az oka, hogy az LLM-ek az MI-nek csak egy részhalmazát alkotják. Ezek a rendszerek nem jelentenek mindenre jó megoldást, és szerinte a jövőben a kisebb, specializáltabb modellek terjednek majd el. Jelenleg minden erőfeszítés és pénz arra irányul, hogy hogyan lehet egyetlen olyan modellt építeni (hatalmas számítási kapacitással), amely majd mindenkinek minden problémáját megoldja. Ezzel szemben áll az, amit prognosztizál: a következő hónapokban, években egyre több olyan specializáltabb modell fog felbukkanni, amelyeket különböző problémák megoldására fejlesztettek ki. Egy banki chatbottal illusztrálja a mondanivalóját. Nincs szükség arra, hogy az ügyfél az élet értelméről beszéljen vele. Egy kisebb, specializáltabb modellt lehet használni erre a feladatra, amely olcsóbb, gyorsabb, és talán vállalat saját infrastruktúráján lehet futtatni. A Hugging Face alapítója elismerte, hogy az LLM-lufi kipukkanása bizonyos mértékig hatással lehet a cégére, de megjegyezte: az MI-ipar már most is annyira nagy és diverzifikált, hogy még ha egy része – például az LLM-ek – túlértékelt is, az nem fogja alapjaiban megrázni sem a teljes területet, sem a vállalkozását.

[Hugging Face CEO says we're in n LLM-bubble, not an AI bubble](#)



Az energiára szomjazó MI-adatközpontok: helyzetkép az USA-ból

A mesterséges intelligencia nagy nyelvi modellek expanzív elterjedésével a közbeszéd fókuszába került az adatközpontok kérdése. Azoknak a „szerver-farmoknak” a témája, amelyek nélkülözhetetlenek az újabb és újabb modellek kifejlesztéséhez (tréningezéséhez), illetve a technológia napi működtetéséhez. Ezek a fontos infrastruktúra elemek elképesztő összegű – dollár milliárdokban mérhető – befektetéseket vonzanak már ma is. Fontos tehát velük kapcsolatban a tisztánlátás, és ezen a téren egy kitűnő, összegző tanulmányt jelentetett meg a napokban a Pew Research. Külön érdekesség, hogy a híres amerikai vállalatot a világ egyik legjelentősebb közvélemény kutató cégeként tartjuk számon, most azonban egy kicsit más terepen is bizonyították szakmai felkészültségüket. Elsősorban természetesen az USA adatközpont hálózatát vették górcső alá, de mivel az MI-technológia (egyik) hazája Amerika, ezért az információk általánosságban is jól hasznosíthatók. A tanulmány olyan fontos alapkérdéseket vesz elő és igyekszik szakszerűen megválaszolni, amelyek szinte mindennaposak ma már az MI-ről szóló globális beszédben, miközben a felbukkanó információk, érvek többnyire az (angolszász fordulattal élve) „anekdotisztikus” kategóriába sorolhatók csupán. A Pew kutatói körbejárták az alapvető fogalmakat: mit tekintünk adatközpontnak; hány adatközpont üzemel ma összesen az USA-ban; mekkora a tényleges energiafelhasználása egy-egy ilyen központnak; honnan származik az ott felhasznált energia; hogyan érinti az adatközpontok működése a környezetet, a természetet és a fogyasztók világát? Természetesen, hivatalos szakmai profiljához kapcsolódva, a Pew alaposan utánanézett annak is, hogyan vélekednek az emberek az adatközpontok gombamód történő szaporodásáról.

[What we know about energy use at US data centers amid the AI boom](#)

