



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

2025 május

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

- Jogkövető MI-ágensek tervezése
- Hogyan visz be hibákat az MI a peres eljárásokba?
- Wikipédia szócikkeken tréningezett MI-modellek

Trendek

- 2025 a humanoid gyári munkás éve lesz
- Példa nélküli az MI fejlődése - trendjelentés
- A mesterséges intelligencia helyzete Kínában
- A kínai Zhipu cég ingyenes MI-ágense országon belül is felerősíti a technológiai versenyt

Működésben

- Terápiás segítő chatbot klinikai vizsgálata: az első eredmények biztatóak
- Afrika, ahol éppen az MI lehet a gyógyír a krónikus tanárhiányra
- Kiábrándulás és mesterséges intelligencia
- Görög tudósok által fejlesztett algoritmus őrködik a homokstrandok tisztaságán
- Mesterséges intelligencia segít megfejteni az állatok „beszédét”





Etika és jog

Jogkövető MI-ágensek tervezése

Az MI-rendszerek fejlődésének új szakaszát a valamilyen fokú önállósággal rendelkező MI-ágensek elterjedése jellemzi, amelyek relatív önállósággal végzik a rájuk bízott feladatokat. Az ezzel kapcsolatos problémák egy része azzal függ össze, hogy az informatikai rendszereket már ma is lehet káros célokra használni (és használják is őket). Az önálló rendszerek megjelenése tovább bonyolítja a helyzetet. Egy önálló MI-rendszer kárt okozhat úgy is, hogy ez nem szándékolt, és úgy is, hogy éppen ez a feladata. A tanulmány témája az, hogy ez utóbbit hogyan lehet megakadályozni. Két alapvető fogalmat tárgyal: egyrészt a jogkövető MI-rendszereket (Law-Following AIs, LFAIs), illetve az MI-bérenceket (AI Henchmen). Az MI-ágenseknek hűségesnek kell lenniük megbízóikhoz, de csak a törvény keretein belül: úgy kell őket megtervezni, hogy megtagadják az illegális cselekmények végrehajtását megbízóik szolgálatában – ezek a jogkövető MI-rendszerek. Velük ellentétben állnak azok az MI-ágensek, amelyeket arra képeztek ki, hogy hűségesen, korlátozások nélkül, a törvényekre való tekintet nélkül kövessék megbízóik kívánságait – ezek az MI-bérencek. Ez a probléma különös élességgel vetődik majd fel közigazgatási, politikai alkalmazások esetén. A problémára számos megoldás létezik már elég hosszú ideje. Ezek azon alapszanak, hogy az automatizált rendszerek viselkedését szabályozó konkrét jogi korlátozásokat úgy építik be az MI-rendszerekbe, hogy a jogi szövegeket formális, géppel olvasható számítógépes kóddá alakítják. A jelenlegi csúcstechnológiájú MI-rendszerek azonban már képesek természetes nyelvű szövegek, köztük törvények olvasására, megértésére és értelmezésére (a szerzők szóhasználatára szerint). Ez a fejlemény új lehetőségeket nyit meg a szabályozásuk terén. Ezen technikai fejlemények alapján azt javasolják a szerzők, hogy az MI-rendszereket olyan módon illesszék be a meglévő emberi jogrendbe, hogy a jogrend hatálya hasonlóan érvényesüljön, ahogyan az emberi magatartást szabályozó törvények körénél is érvényesül. Ezzel a megoldással az MI-ágensek jogképesek lennének, jogi kötelezettség terhelné őket. Ez ugyan alapos változtatás lenne a jog jelenlegi rendszerében, de nem példa nélküli. A szerzők a jogi személy fogalmára utalnak, és arra, hogy egy adott vállalatot is terhelhetnek kötelezettségek. A szerzők álláspontja szerint ez a javaslat összhangban van több vezető MI-vállalat kialakulóban lévő biztonsági gyakorlatával is.

[Law-Following AI: Designing AI Agents to Obey Human Laws](#)

Hogyan visz be hibákat az MI a peres eljárásokba?

Mindenki használja már a chatbotokat, nincs ez másként a jog világában sem. Ez persze azt is jelenti, hogy meg kellene tanulni együtt élni a jól ismert hibáikkal, megfelelő stratégiákat kialakítani ezek kezelésére. Az utóbbi hetekben számos amerikai hír jelezte, hogy ennek a folyamatnak egyelőre csak az elején tartunk. Egyre gyakrabban jelennek meg a jogi eljárások beadványaiban az MI-rendszerek által létrehozott hallucinációk, és ezekben az iratokban nem létező példákra és törvényhelyekre hivatkoznak. És mindez már elkezdte zavarni a bírókat. Összességében ezek az esetek egy komoly problémára utalnak. A bíróságok olyan dokumentumokra támaszkodnak, amelyek pontosak és hivatkozásokkal alátámasztottak – két tulajdonság, amelyet az MI-modellek, annak ellenére, hogy az időt megtakarítani vágyó ügyvédek alkalmazták őket, gyakran nem tudnak biztosítani. Ezeket a hibákat az esetek többségében még ki szokták szűrni, de nem nehéz elképzelni, hogy hamarosan akár egy bíró döntését is befolyásolhatja majd olyan tényező, amelyet teljes egészében az MI talált ki, és senki sem fogja észrevenni. Nem ismeretlen cégekről, elszórt egyszeri esetekről van szó. Nagynevű ügyvédi irodák kínos pillanatai ezek. Az a tény, hogy tekintélyes jogászok, akik szakmájuk szerint szakértői annak, hogy érvelési hibákat vegyenek észre, ismételten átsiklanak az MI-rendszerek hallucináció felett, csak részben magyarázható időhiánnyal (vagy más hasonló okkal). Úgy tűnik, hogy a nyelvi modelleket sokan eleve tekintélyként kezelik, hiszen bonyolult kérdésekre átgondolt, intelligensnek tűnő választ adnak. Az ügyvédek megszokták, hogy ellenőrzik a fiatal ügyvédek és gyakornokok munkáját, de úgy tűnik, hogy vannak, akik ezt a szkeptikus hozzáállást nem alkalmazzák az MI-rendszerek által adott válaszok esetében. A probléma, a hallucináció, mint a generatív MI-rendszerek alapvető, enyhíthető, de meg nem szüntethető tulajdonsága már a ChatGPT közel három évvel ezelőtti megjelenése óta ismert, de az ajánlott megoldás azóta nem sokat változott: ne higgy el mindent, amit olvasol, és ellenőrizd, amit az MI-modell mond neked. Mivel egyre inkább az MI-eszközöktől eredő információval van dolgunk, ezért az MI-modellek egyre több különböző eszközbe kerülnek, emiatt kétséges, hogy mennyire reális tömegmérétekből ez a stratégia. Nehezíti a helyzetet, hogy számos nagy szolgáltató úgy árulja a jogászoknak szánt generatív MI-n alapuló termékeit, hogy ezek megbízhatóan pontosak.

[How AI is introducing errors into courtrooms](#)

Wikipédia szócikkeken tréningezett MI-modellek

Az exponenciálisan növekvő ütemben fejlesztett és piacra dobott nagy nyelvi modellekkel kapcsolatban már jó ideje ismert az egyik legkomolyabb korlát: az új modellek betanításához és a már meglevő modellek továbbképzéséhez, finomhangolásához szükséges gigantikus méretű szövegmenyiség végessége. Már egy jó ideje tartja magát az álláspont, hogy a még „szűz”, tehát az LLM-ek tréningezéséhez eddig fel nem használt szövegtérlet gyakorlatilag kifogyott. Ennek a ténynek a folyamata egy sor fejlesztési törekvés – az új modellek szintetikus adatkészleteken történő betanítása, illetve a már meglevő modellek „tutorként” történő felhasználása, azaz tudáskészletének „lehívása” –, és nemkülönben számos nagyon komoly következménnyel kecsegtető jogi eljárás, szerzői jogi per. Ebben az összefüggésben különösen érdekes, előremutató a Wikipédia Enterprise cég egy mostani kezdeményezése. Az LLM-fejlesztők eddig is előszeretettel használták modelljeik betanításához a Wikipédia szócikkeit – a platform tulajdonosának megkérdezése és beleegyezése nélkül. Most a Wikipédia cég (a jól ismert mondást alkalmazva: „Ha valamit nem tudsz megakadályozni,

akkor állj annak az élére!”) olyan strukturált adatkészleteket állít elő a Wikipédia angol és francia nyelvű szócikkeiből, amelyek kifejezetten a nagy nyelvi modellek betanításának igényeit tartják szem előtt. Azaz, a nyers szóhalmazokat eredményező eddigi „webgereblyezés” (scraping) helyett a fejlesztők mostantól tisztított, gépi olvasásra előkészített adatbázisokhoz juthatnak hozzá.

[Training AI Models on Wikipedia Content](#)





Trendek

2025 a humanoid gyári munkás éve lesz

Úgy tűnik, hogy 2025 lesz az év, amikor a többcélú humanoid robotok, amelyek eddig főként kutatólaboratóriumokban voltak használatosak, kereskedelmi forgalomba kerülnek. A Goldman Sachs 2024-es jelentése szerint a humanoid robotok piaca 2035-re 38 milliárd dollárra nőhet, ami több mint hatszorosa a cég egy évvel korábbi előrejelzésének. A humanoid robotok alapvető ígérete, hogy képesek lesznek több feladat között váltani, akár csak emberi társaik. Ez alapvetően eltérő megközelítés a hagyományos futószalagos automatizálástól, amely a gyártáshoz szükséges specifikus feladatok köré építi fel a teljes környezetet. „A célra szabott automatizálási megoldás mindig nagyobb teljesítményt és alacsonyabb költségeket jelent az adott cél tekintetében” – mondják a szakértők. Ez nagyon jó, ha az adott feladatot 24 órában, a hét minden napjában végzik. De olyan feladatok esetében, amelyeket nem kell éjjel-nappal végezni, egy rugalmas robot lehet a termelékenyebb megoldás. Másképpen is lehet fogalmazni: a gyárat eleve úgy alakítják ki, hogy a robotok biztonságosan mozogjanak. A következő nagy lépés azonban az olyan robotok létrehozása lesz, amelyek az emberek világában, a gyorsan változó körülmények között is (mindkét félre tekintettel) biztonságosan működnek. Az egyik legnagyobb akadálya annak, hogy a többcélú humanoidok valóban hasznosak legyenek, hogy rengeteg időre van szükség ahhoz, hogy a robotokat megtanítsák a különböző feladatok elvégzésére. A szakértők azonban úgy vélik, hogy a nagy nyelvi modellek által elért hatalmas fejlődés segíthet megoldani ezt a problémát. Ez is az egyik oka annak, hogy humanoid robotprojektek jelennek meg olyan vállalatoknál, mint az Apple, a Meta és a Tesla, amelyek máris virágzó MI-laboratóriumokkal rendelkeznek. Miután a robotok képesek lesznek a készségek gyors elsajátítására, jobban fel lesznek készítve arra, hogy az emberi munkásokhoz hasonlóan feladatról feladatra váltsanak. A természetes nyelvfeldolgozás terén elért folyamatos fejlődés azt is jelenti, hogy a robotok hangutasítások alapján is irányíthatók lesznek. De még mindig vannak biztonsági aggályok a robotok és az emberek együttes munkavégzésével kapcsolatban, és sok szélsőséges esetet kell majd kezelni. A szakértő szerint a legfőbb kihívás a megbízhatóság biztosítása lesz.

[2025 Is the Year of the Humanoid Robot Factory Worker](#)

Példa nélküli az MI fejlődése - trendjelentés

Mary Meeker a híres kockázati tőkebefektető (egyben kommentátor is) felújította trendjelentései megszokott sorozatát. Frissen megjelent jelentésében 340 dián keresztül bizonyítja, hogy a digitális világ mostani viharos eseményei, amelyet az MI-rendszerek egészen látványos felívelésének köszönhetünk, mások az eddig tapasztalt „forradalmaktól” – pl. a mobil távközlés, a közösségi média, a felhőinformatika eseményeitől. Úgy is fogalmazhatunk, hogy azért gondoljuk azt, hogy most példátlan dolgok történnek, mert most tényleg példátlan dolgok történnek. A 340 oldalas, adatokban hihetetlenül gazdag, az MI-fejlesztések szinte minden vonatkozását tárgyaló dokumentum 51 oldalán fordul elő a „példátlan” kifejezés. A ChatGPT 17 hónap alatt elérte a 800 millió felhasználót: ez példátlan. A fejlesztésekkel foglalkozó vállalatok száma és az a sebesség, amellyel sok vállalat is magas éves ismétlődő bevételi rátákat ér el (azaz szolgáltatásaikat rendszeresen használják más vállalatok): szintén példátlan. A használati költségek csökkenésének sebessége: példátlan. Míg a modellek képzésének költségei (amelyek egyébként szintén példátlanul magasak) akár 1 milliárd dollárra is rúghatnak, a következtetés (inference, azaz a tényleges használat) költségei – például azok számára, akik fizetnek a technológia használatáért – már 99%-kal csökkentek két év alatt, ha egy millió tokenre vetítve számoljuk a költségeket. Az az ütem, ahogy a versenytársak egymás funkcióit utánozzák, az eredeti fejlesztési költségek töredékéért, beleértve a nyílt forráskódú opciókat, különösen a kínai modellek esetében: példátlan. Az egyetlen terület, ahol a mesterséges intelligencia nem előzött meg minden más technológiai forradalmat, az a pénzügyi megtérülés. Miközben a kockázati tőkebefektetők olyan gyorsan öntik a pénzt az MI-vállalkozásokba, amilyen gyorsan csak tudják, az MI-vállalatok és a felhőszolgáltatók is gyorsan „égetik” a rendelkezésükre álló készpénzt. A mesterséges intelligencia hatalmas infrastrukturális beruházásokat igényel. Ez a példátlan ütem jó a fogyasztóknak és a vállalkozásoknak, a gyors fejlesztések haszonélvezőinek, miközben a verseny csökkenti a költségeket. De még nem lehet tudni, hogy a jelenlegi vállalatok közül melyik lesz hosszú távon nyereséges, a következő generációs technológiai óriás. Az egész MI-világ azzal az aggodalommal kénytelen megküzdeni, hogy a mérhetetlen várakozások által tüzelt folyamat előbb-utóbb váratlan véget ér, szembesülve a realitásokkal. A szerző álláspontja határozottan optimista. Szerinte hasonló a helyzet, mint az e-mail és webes keresés korai napjaiban, ezek is olyan technológiák, amelyek alapvetően megváltoztatták a világunkat. A mesterséges intelligencia jobb/gyorsabb/olcsóbb hatásai ugyanolyan varázslatosnak tűnnek, csak még sokkal gyorsabban hatnak. Az optimista befektető alapállását idézi jóváhagyólag: „Statistikailag nézve, a világ nem ér véget olyan gyakran.”

[Trends – Artificial Intelligence](#)

A mesterséges intelligencia helyzete Kínában

Az Artificial Analysis cég egy független tanácsadó vállalkozás, amely az MI-rendszerek teljesítőképességének összevető mérésével és iparági tanácsadással foglalkozik. E tevékenységük keretében jelentették meg összefoglalójukat a kínai MI-ipar aktuális helyzetéről. A teljes jelentés csak előfizetőknek érhető el, a továbbiakban a nyilvános összefoglalót ismertetjük – ez is meglehetősen adatgazdag és jól alátámasztott. A jelentés mondanivalójának ismertetése előtt két terminológiai megjegyzés. A jelentésben a „frontier AI”-ról van szó, ezt úttörő, élenjáró MI-nek fordítjuk. Ezek az MI-vel elérhető eredmények határait feszegető, legújabb generációs MI-rendszerek. A másik fontos fogalom, amiről szó lesz, az a „súly” (weight). Ez a rendszer alkotóelemeihez (neuronok) kapcsolt számszerű információ, amely végső soron a tanításukhoz használt adatok viszonyát reprezentálja. A súly a rendszer „tudása”. Kína vezető mesterséges intelligencia laboratóriumai most közelebb vannak a vezető amerikai fejlesztőkhöz, mint valaha. Hátrányuk több mint egy év volt, ez ma már kevesebb mint három hónap. Elsősorban a DeepSeek és az Alibaba a kínai élenjáró fejlesztések meghatározó cégei. A kínai élenjáró, nyílt súlyú fejlesztések 2024 novemberében haladták meg az Egyesült Államokét, miután az Alibaba kiadta a QwQ 32B Preview-t, az R1 pedig megszilárdította ezt az előnyt. A kínai MI-laboratóriumok nyílt súlyok terén betöltött vezető szerepe tükrözi a kínai MI-laboratóriumok azon gyakorlatát, hogy gyakran közzéteszik vezető modelljeik súlyait - ellentétben az amerikai gyakorlattal. A vezető kínai mesterséges intelligencia laboratóriumok, a DeepSeek és az Alibaba folyamatosan új modelleket hoznak forgalomba, 2024 végén a DeepSeek vette át a vezetést. A DeepSeek R1 (2025. január) volt az első nyílt súlyozási érvelési modell, amely versenyképes volt az OpenAI o1-ével. A DeepSeek R1 0528 (2025. május) a jelenleg elérhető legintelligensebb nyílt súlyozási modell, és továbbra is intelligenciaelőnyt élvez az Alibaba Qwen3 235B A22B modelljével szemben, amely a kínai MI laboratóriumok vezető modellje. Mindkét vállalat nyitott súlyozási stratégiát alkalmaz, támogatva modelljeik széles körű elterjedését belföldön, valamint nemzetközi szinten is. A nyílt rendszerek alkalmazása fontos stratégiai eszköz a kínai felzárkózásban, mert ezek a rendszerek a rugalmasságuknál és olcsóságuknál fogva, összekapcsolódva a kiemelkedően fejlett ipari ökoszisztémával, az MI-alapmodellekre épülő alkalmazások széles köre előtt nyitják meg az utat. A jelentésben sokrétű, bőséges információt kapunk a kínai MI-ökoszisztémáról, a főbb szereplőket számos jellemző alapján tárgyalva. További erőssége az anyagnak, hogy folyamatosan összevetik a kínai szereplőket az amerikai nagy MI-fejlesztőkkel.

[Q2 2025 State of AI – China Report](#)

A kínai Zhipu cég ingyenes MI-ágense országon belül is felerősíti a technológiai versenyt

A „mesterséges intelligencia versenyfutás” ma már az emergens technológiákkal foglalkozó műszaki szakirodalomban is éppúgy gyakori fordulat, mint a technológia és a globális hatalom összefüggéseit boncolgató techno-geopolitikai írásokban. Ebben az összefüggésben általában a két szuperhatalom, az Egyesült Államok és Kína technológiai elsőségért való vetélkedéséről van szó. Habár Kína – részben az ismereteink korlátozott voltából fakadó módon, részben az erős propaganda pengeváltások miatt is – egy monolitikus tömbként jelenik meg (társadalmát, gazdaságát, politikai szisztémáját, és természetesen technológiai szcénáját tekintve is), a valóság ennél sokkal árnyaltabb, sokszínűbb. A Reuters egy híradása éppen erre irányítja a figyelmet: egy általunk alig hallott nevű kínai technológiai

szereplő, a Zhipu AI cég is berobbant az MI-ágensek egyre pezsgőbb fejlesztői közösségébe. Ráadásul bizonyos értelemben igazodva a jelentős kínai MI-cégek filozófiájához (gondoljunk csak a DeepSeek, vagy akár a Qwen néhány hónappal ezelőtti politikájára), ingyenesen téve elérhetővé a fejlesztők, vagy akár a laikus felhasználók számára fejlett modelljüket. A lépés jelentőségét segít megérteni, ha felidézzük, hogy a nemzetközi szakmai közegekben szinte sokkot kiváltó másik „ismeretlen” kínai ágens-fejlesztő, a világ első (csaknem) általános célú ágensével előrukkoló Manus jelentős havi díjért (199 dollárért) teszi hozzáférhetővé modelljét. A szakértők arra hívják fel a figyelmet, hogy a Zhipu lépése jelentősen felbolygathatja a kínai belső MI-fejlesztői piacot: azaz, az MI-versenyfutás a technológiai elsőségért immár nem csupán a két nagyhatalom között zajlik, de érezhetően erősödik Kína belső piacán is.

[China's Zhipu AI launches free AI agent, intensifying domestic tech race](#)





Működésben

Terápiás segítő chatbot klinikai vizsgálata: az első eredmények biztatóak

A mesterséges intelligencia technológiák, ezen belül is a nagy nyelvi modellek egyik kiemelt alkalmazási területe (legalábbis a világ néhány – jellemzően fejlett – társadalmában) a „terápiás alkalmazás”. Ebben az összefüggésben a „terápia” nem általában vett orvosi kezelést takar, hanem kifejezetten a lelki problémákkal való, szoftver alapú foglalkozást. A „posztmodern” társadalmak szorongásokkal, neurózisokkal terhelt embere számára ez az új technológia egyfajta tökéletes megoldást kínál – vagy legalábbis ígér. Az érzelmi tartalmakat (pontosabban: a szöveges tartalmak érzelmi töltéseit) kiválóan érzékelő LLM-ek kézenfekvő megoldást kínálnak a személyes lelki problémáikra megoldást kereső vagy azokról legalább valakivel beszélgetni vágyó ember számára. Nyilvánvaló előny az is, hogy a „terápiás nagy nyelvi modell” mindenhol és a nap bármely pillanatában kéznél van. És akkor a pszichológiai, vagy pszichiátriai konzultációk jelentős költségeiről még nem is beszéltünk. A felhasználási terület ugyanakkor már egy jó ideje az MI-technológiával kapcsolatos etikai viták középpontjában áll. Nyilván könnyű belátni, hogy miért. A sok szempontból fekete dobozként működő, illetve téves vagy „fabrikált” válaszadásra (más szóval hallucinációra) is hajlamos MI-modellek adott esetben nagyobb egészségügyi kárt okozhatnak a hozzájuk forduló, bennük megbízó emberekben. A társadalmi igény ugyanakkor, ahogy említettük, jelentős, ezért komoly áttörésnek tekinthető a Dartmouth Egyetem által kifejlesztett terápiás LLM debütálása. A Therabot névre keresztelt szoftver már túl van az első klinikai próbákon is. Az általános célú chatbot összességében statisztikailag mérhető, jelentős javulást hozott a kísérletbe bevont páciensek tüneteinek tekintetében.

[First Therapy Chatbot Trial Yields Mental Health Benefits](#)

Afrika, ahol éppen az MI lehet a gyógyír a krónikus tanárhiányra

Afrika, a Globális Délnek nevezett harmadik világ egyik leginkább egyenetlen fejlettségű régiója. A kontinenst a gyarmati időszak vége óta sújtó problémák, a minden szempontból strukturálisnak nevezhető elmaradottság alig akar búcsút inteni ennek a vidéknek. A problémák egyik legkritikusabb területe az oktatás; éppen az a társadalmi alrendszer, amely tömegeknek, talán az egész kontinensnek kínálhatná – belátható időn belül – a felemelkedés, felzárkózás ígérését. Az ezen a területen megragadható problémák egyre erősebben kötődnek a földrész (elsősorban az egyenlítői Afrika országai) tanárhiányához. Kenya például mintegy 100 ezer betöltetlen pedagógusi állás mellett próbálja legalább szinten tartani közoktatásának színvonalát. A túlszűfolt termék, nagyméretű osztályok mindennapiak, különösen a főváros, Nairobi körzetében. A túlterheltségtől hamar kiégő oktatói gárda tovább súlyosbítja a már így is kritikus tanárhiányt. Az egyre borúsabb képen a laikusok számára is hozzáférhető nagy nyelvi modellek megjelenése hozott némi derengő fényt. 2023-tól kezdve, a néhány hónappal korábban elérhetővé tett ChatGPT-t felhasználva egy helyi civil szervezet speciális technológiai továbbképző tanfolyamokat szervezett: az órai tananyagok tervezésétől kezdve a diákok írásbeli munkájának ellenőrzéséig, a javítás sok időt elvevő tevékenységéig számtalan

területen támaszkodhatnak az új eszközökre a helyi pedagógusok. Az MI-technológia segítő lehetőségei azonban korlátozottak maradnak: Kenyában ugyanis, Afrika számos országához hasonlóan messze nem általános a megbízható minőségű, folyamatos internet kapcsolat. Egy közelmúltbeli felmérés (amely a térség 10 országának konnektivitását igyekezett feltérképezni) szerint a térség országaiban az iskoláknak átlagosan legfeljebb 35%-a képes megfelelő minőségben csatlakozni a világhálózathoz. A Kenyában MI-képzéseket szolgáltató civil szervezet, a Kalasik saját chatbotot fejlesztett (finomhangolt) egy korábbi OpenAI modellre alapozva. Bár a szoftverrel alapvetően elégedettek a pilot program keretében azt használó tanárok, az alapvetően a kenyai hivatalos kerettantervhez igazodó MI néha „elkalandozik” az interneten. Néhány kenyai start-up cég más MI-alapú eszközök fejlesztésével igyekszik segíteni a közoktatásban dolgozó tanárok munkáját: a Nyansapo AI például olyan beszédfelismerő szoftverrel állt elő, amely szintén egy OpenAI modell finomhangolásával biztosította, hogy az applikáció felismerje a kenyai angol dialektust.

[Kenya's solution to teacher shortage: Embrace AI](#)

Kiábrándulás és mesterséges intelligencia

Sok vállalat számára a generatív mesterséges intelligencia ígéretével kapcsolatos izgalom helyét a technológia produktív felhasználásának nehézségeivel kapcsolatos bosszúság vette át. Egyes felmérések szerint a generatív MI-pilot projekteik nagy részét feladó vállalatok aránya 42%-ra emelkedett, szemben a tavalyi 17%-kal. A fogyasztók viszont továbbra is lelkesen fogadják a generatív MI-t. Sam Altman, az OpenAI vezetője nemrégiben elmondta, hogy a ChatGPT chatbotot hetente mintegy 800 millió ember használja, ami kétszerese a februári adatnak. Néhányan már rendszeresen alkalmazzák a technológiát a munkájuk során is. A generatív MI átalakító potenciálja azonban csak akkor valósulhat meg, ha a vállalatok széles köre szisztematikusan beépíti termékeibe és működésébe. A lassú előrehaladás miatt sok vezető csúszik bele a „kiábrándulás völgyébe” – mondja John Lovelock, a Gartner munkatársa, utalva a tanácsadó cég híres „hype cycle” modelljének azon szakaszára, amely az új technológia által kiváltott eufória után következik be. A vállalatok számára számos ok miatt küzdelmes a generatív MI rendszerszerű használata. Adataik gyakran egymással nem kommunikáló, elszigetelt rendszerekben vannak, amelyeket elavult IT-rendszerek működtetnek. Nem könnyű hozzáértő szakembert találni sem. Bármennyire is nyilvánvaló a technológia átalakító potenciálja, minimalizálniuk kell annak kockázatát, hogy az át nem gondolt bevezetés miatt a vállalat jó hírre se sérüljön. Az MI-rendszerek alkalmazásának sikereit ritkán szokták árbevételben vagy elért nyereségben megadni. Jelenleg azok a nagy felhőszolgáltatók, amelyek az MI-rendszerek széles körű vállalati alkalmazásához szükséges hardveres és szoftveres feltételeket szolgáltatnák, mint például az Alphabet, az Amazon és a Microsoft, csekély bevételekről számolnak be (legalábbis befektetéseik nagyságához képest), ezek is főként az MI-rendszereket fejlesztő cégektől és induló vállalkozásoktól származnak – amelyek egy részét maguk ezek az óriáscégek finanszírozzák. Sok vállalat szerint nem kifinomultabb (okosabb) MI-modellekre van szükségük, hanem gazdagabb, bővebb hasznosítási eljárásokra, módszerekre, hogy a technológiát hasznosíthassák. Jelenleg úgy tűnik, hogy a ténylegesen használható képességek és a vállalati igények között nincs összhang.

[Welcome to the AI trough of disillusionment](#)

Görög tudósok által fejlesztett algoritmus őrködik a homokstrandok tisztaságán

Az MI-fejlesztéseket monitorozó munkánk során igyekszünk beszámolni az MI-alkalmazások terén felbukkanó „színes” fejleményekről is. Nem csak arról van itt szó, hogy egy-egy könnyedebb hírrel segítsük az olvasó figyelmének ébren tartását. Meggyőződésünk, hogy a technológiák esetében nem maga az innováció az igazi lényeg, hanem a társadalmi adaptáció, azaz hogy a társadalom tagjai – lehetőleg minél többen és egyre újabb területeken – fedezzék fel az új technológiában rejlő előnyöket. Nem a találmány, hanem annak az alkalmazása (vagy annak hiánya) az igazi kulcs. Ebben a fontos, társadalmi adaptációs folyamatban a nagy, meghatározó kulcsterületeken történő alkalmazásokkal egyenértékűnek tartjuk a nem életbevágóan fontos területeken való alkalmazásokat is. Az emberi társadalmak élete valójában csekély jelentőségű, de mégis sokat számító alrendszerekből, tevékenységekből, folyamatokból kerekedik ki. Egy jellegzetesen ide kapcsolódó fejlesztési hír Görögországba kalauzol minket. Kutatók egy most fejlesztett algoritmus segítségével igyekeznek gyógyírt találni a görög tengerpart tipikus fenyegetettségére: a tömegturizmus környezetpusztító hatásának mérséklésére. A Nyugat- Peloponnészosz 12 népszerű tengerparti strandját felölelő projekt egy sor környezeti tényezőt vet össze a partszakaszt érő embertömeg terhelésével. A csapat felmérte az egyes strandok „teherbíró képességét”, vagyis azt, hogy hány látogatót tudnak befogadni anélkül, hogy károsítanak az ökoszisztémát. Az eredmények szerint egyes strandok a fenntartható szintnél kétszer-háromszor több látogatót fogadnak. Az algoritmus adat alapú megközelítéssel igyekszik megoldást kínálni arra, hogy a lehető legtöbb látogató élvezhesse a tájat, egyúttal megőrizve a környezet épségét is.

[Greek scientists develop algorithm to help keep beaches pristine](#)

Mesterséges intelligencia segít megfejteni az állatok „beszédét”

A mesterséges intelligencia technológiák szívós tempóban terjeszkednek az élet különböző területeire, ám az esetek többségében megmaradnak a mi fajunk, tehát a homo sapiens egyéni és közösségi életének színtereiben. Van azonban néhány érdekes kísérlet arra, hogy ezeket az új technológiákat merőben forradalmi eszközként használva egy – ősi idők óta vágyott – hidat építsenek az ember és az emberen kívüli világ között. A kezdetek, meglepő módon, legalább egy fél évszázadnyi távolságra nyúlnak vissza. A XX. század közepe táján egész sor (az akkori korok elektroakusztikai technológiáit is felhasználó) kísérlet indult például a púposhátú bálnák kommunikációjának megfejtésére. Az emberszabásúak természetesen külön figyelmet ébresztettek mindig is: etológusok egész generációja igyekezett a csimpánzok, gorillák hang- és jelbeszédét megfejteni. A mesterséges intelligencia technológiák most kibontakozó forradalma talán végre áttörést hozhat ezen a téren is. A gépi tanulás és a robotika fejlődése lehetővé teszi, hogy soha nem látott mélységben és pontossággal tanulmányozzuk az állatok kommunikációját. A tudósok ma már rendelkeznek azokkal az eszközökkel és képességekkel, amelyekkel összeállíthatják és elemezhetik a spekulációkon túllépő, fordításra és akár a fajok közötti kommunikációra is alkalmas hatalmas adatállományokat. Azonban ahogyan a mesterséges intelligencia és más csúcstechnológiai eszközök soha nem látott lehetőségeket teremtenek az állatok megértéséhez és védelméhez, úgy komoly kockázatot is jelentenek a további manipuláció és kizsákmányolás szempontjából. Jogi és etikai korlátok nélkül azt kockáztatjuk, hogy ezeknek a technológiáknak a használatával árthatunk is azoknak az élőlényeknek, amelyeket megérteni szeretnénk. Az állatok jogainak és jólétének védelmére törekvő tudósoknak kellő gondossággal és óvatossággal kell közelíteniük ehhez a munkához, az állatok jólétét előtérbe helyezve a felfedezés izgalmával szemben.

[Researchers are Using AI to Understand what Animals are Saying](#)

