



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

2025 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

- Ismerkedjünk meg az MI-t alkalmazó, kísérletező kedvű bírókkal
- Az ördög és az angyal között: a mesterséges intelligenciával való együttélés kockázatai
- A tudományos könyvkiadás etalonjának számító Springer Nature gépi tanulásról szóló új könyve – tele van nem létező hivatkozásokkal

Trendek

- Három lehetséges MI-jövő
- Kína mesterséges intelligencia jövőképe más – és lehet, hogy okosabb
- MI adatközpontokat „zöldít” egy akkumulátor újrahasznosító üzem
- Vita a mesterséges általános intelligencia (AGI) örvén

Működésben

- A generatív MI-vel támogatott programozás rejtett költségei
- Algoritmikus segédedző dolgozik az Országos Amerikaifutball Liga (NFL) trénereinek keze alá
- MI-alapú megfigyelő rendszer védelmezi az Indonézia körüli veszélyes tengerek biztonságát
- A ChatGPT legutóbbi modellváltásával eltűnő régi verziót gyászoló felhasználók





Etika és jog

Ismerkedjünk meg az MI-t alkalmazó, kísérletező kedvű bírókkal

Széles körben elterjedt az amerikai igazságszolgáltatásban az MI-rendszerek használata. Mind az ügyvédek, mind a bírók használják, azonban a bírói alkalmazási gyakorlatok számos problémát hoztak a felszínre. A bírók bízhatnak abban, hogy a megfelelő óvintézkedésekkel a technológia felgyorsíthatja a peres eljáráshoz kapcsolódó jogi kutatásokat, összefoglalhatja az ügyeket, rutinszerű végzéseket fogalmazhat meg, és összességében segíthet felgyorsítani a bírósági rendszer működését, mert a bíróságok az Egyesült Államok számos részén súlyos ügghátralékkal küszködnek. Ezen korai alkalmazói kísérletek eredményei két dolgot tesznek világossá. Egyrészt, a rutinfeladatok kategóriáját – amelyekben a mesterséges intelligencia emberi ítélőképesség nélkül is segíthet – nehéz meghatározni. Bíróról bíróra és bírósági ügyről bírósági ügyre változik, hogy az MI-rendszerek és az emberek számára alkalmas feladatok közötti határt hol húzzák meg. Számos olyan eset is van, amelyekben az MI-rendszerek nem alkalmazhatók jó eredménnyel, kellő gondosság hiányában. Ha például arra kérjük a rendszert, hogy összegezzon egy nagyméretű dokumentumot, az eredmény drasztikusan eltérő lehet attól függően, hogy a modellt arra képezték-e ki, hogy általános közönségnek vagy jogi célközönségnek összegezzon a dokumentumokat. Az MI-rendszerek a logikán alapuló feladatokkal is küzdenek, valamint például egy eset eseményeinek sorrendbe állításával is. Másrészt, míg az ügyvédek szigorú ellenőrzéssel szembesülnek a peres eljárás során, és ennek során kiderülnek az MI-rendszerek naiv (nem ellenőrzött) használatából eredő hibák, a bírák esetében – a peres folyamat sajátosságai miatt – sokkal nehezebb felfedezni a hibákat, mielőtt azok kárt okoznának. Született néhány olyan ítélet, amelyet kínos, a használt MI-rendszereknek köszönhető hibák jellemeztek. Ezek a hibák alááshatják a közvélemény bizalmát a bíróságok legitimitásában, krízishelyzetet okozhat ez a probléma. A mesterséges intelligencia bizonyos szűk és ellenőrzött alkalmazásai – tanúvallomások összefoglalása, gyors írásbeli visszajelzések előállítása – időt takaríthatnak meg, és jó eredményeket hozhatnak, ha a bírák úgy kezelik az MI-rendszerek által előállított információt, mint egy kezdő gyakornokét, alaposan ellenőrizve annak pontosságát. De a bírói munka nagy részét úgy kell elképzelni, mint ami egy tiszta, fehér lappal indul: amikor egy összetett ügyben kell ítélezni, kezdetben csupán egy üres lap van a bíró előtt, azt kell megtölteni tartalommal, és ehhez nehéz döntéseket kell hozni. A szakértő szerint ezeknek a döntéseknek a végiggondolása valójában a bírói munka lényege. Amennyiben az első vázlatot egy MI-rendszerrel készítik el, akkor ez pontosan ezt a lényegi részt sodorja veszélybe. Az nem igazságszolgáltatás, ha egy gyermekelhelyezési perben a bíró helyett egy nagy nyelvi rendszer hozza meg a döntést.

[Meet the early-adopter judges using AI](#)

Az ördög és az angyal között: a mesterséges intelligenciával való együttélés kockázatai

A mesterséges intelligencia fejlődését szélsőséges vélemények kísérik. Az egyik véglet démonizálja, úgy állítja be, hogy nyakunkon a végétélet és a világvége az elszabaduló és önálló életre kelő MI-rendszerek miatt. A másik, ezzel ellentétes szemlélet túlzott bizalommal tekint a fejlemények elé, azt gondolva, hogy az MI-rendszerek működése probléma nélkül emberi ellenőrzés alatt tartható. Ezekkel a nézetekkel szemben a jelentős szakmai folyóirat, az AI & Society egyik friss cikkének szerzői bevezetik a „szimbiotikus kockázat” fogalmát. A szimbiózis fogalma a biológiából származik, két vagy több különböző faj egyedeinek (általában egymásra utalt) szoros együttélését jelenti. Norbert Wiener nyomán hangsúlyozzák, hogy az MI soha nem pusztán passzív eszköz vagy egyszerű fenyegetés, hanem egy összetett rendszer, amely az emberi társadalommal való kölcsönös interakció és együttes-evolúció révén jön létre és működik. A kockázatok nem egyszerűen magából a technológiából fakadnak, és nem is kizárólag emberi hibákból erednek; a szimbiotikus interakcióban történő folyamatos generálás, felhalmozódás és felerősödés eredményei. Ezért a szimbiotikus kockázat különbözik a gépek meghibásodásától vagy a technikai elfogultságtól; olyan mélyen gyökerező veszély, amely a mindennapi ember-MI interakcióban a túlzott bizalom és függőség révén keletkezik. Mivel társadalmi probléma, ezért össztársadalmi, interdiszciplináris együttműködést igényel a kezelése. Az adatok átláthatóságát és az algoritmikus magyarázhatóságot prioritásként kell kezelni a valós kockázatok, például az algoritmikus torzítás, a „visszhangkamrák” és az adatokkal való visszaélés mérséklése érdekében. Ugyanakkor meg kell erősíteni a közoktatást, hogy segítsen az embereknek jobban megérteni a mesterséges intelligencia korlátait és potenciális kockázatait, valamint racionális társadalmi elvárásokat kialakítani a mesterséges intelligenciával szemben. Az MI démonizálása és istenítése közötti ingadozás elkerüléséhez túl kell lépni a „kontroll” kontra „kontrollvesztés” leegyszerűsített elképzelésein. A mesterséges intelligencia összetett társadalmi-technikai rendszer, ami azt jelenti, hogy a bizonytalanság kiküszöbölhetetlen. A szimbiotikus kockázatokkal való hatékony szembenézés a társadalom sokszínűségét tükröző együttműködéssel valósítható meg.

[Between the devil and the angel: confronting the symbolic risks of artificial intelligence](#)

A tudományos könyvkiadás etalonjának számító Springer Nature gépi tanulásról szóló új könyve – tele van nem létező hivatkozásokkal

Az internetet az elmúlt hónapokban szinte elöntötték a komolynak látszó, ám valójában MI-chatbotok által fabrikált publikációk. Ezek jobbára az angolszász megfogalmazással „fikciós” könyveknek nevezett „szépirodalmi alkotások”, ám ezek minőségére vonatkozóan bőségesen indokolt itt az idézőjelek használata. Ezzel párhuzamosan a „nem-fikciós” irodalom előállításában is kezdenek jeleskedni az MI-t használó „szerzők”, de összességében ez a jelenség mégis megmaradt a csináld magad könyvek (egyáltalán nem lebecsülendő méretű) piacán. Legalábbis mindeddig. A közelmúltban feltűntek ugyanis olyan példák, amelyek arra utalnak, hogy az MI-vel „íratott” fércmunkák utat találnak a komoly szakmai kiadók gyakorlatába, illetve kínálatába is. A részben a szakmai hamisítások kiszűrésére specializálódott Retraction Watch a közelmúltban publikált egy riasztó esettanulmányt, amely kimutatta: az ilyen géppel készített áltudományos kiadványok most már a világ olyan elismert tudományos szakkiadóinak a kínálatát is megfertőzik, mint amilyen a Springer Nature. Az általuk bemutatott példa ráadásul egy gépi tanulással foglalkozó szakkönyvnél mutat ki kitalált, fabrikált hivatkozásokat. Ez – ma már közismerten – a nyelvi modellek által „generált” szövegek egyik árukodó jellegzetessége. A Retraction Watch emberei felvették a kapcsolatot számos, a vizsgált könyvben

említett hivatkozás szerzőivel, akik egyértelműen megerősítették, hogy a Springer Nature-nál megjelent szakmunkában feltüntetett bibliográfiai hivatkozások tévesek, pontatlanok, vagy egyenesen kitalációk. Egyesekben felmerülhet a félelem: az MI-technológiák terjedése immár nem csak a felsőoktatást forgathatja fel gyökerestül, de a tudomány egészét fenyegetheti.

[Springer Nature book on machine learning is full of made-up citations](#)





Trendek

Három lehetséges MI-jövő

Egy nemrégiben megjelent tanulmány szerzője Charles Dickenstől vett idézettel írja le a mai helyzetet, az MI körüli paradoxonok világát (egy jövőbeli történész szemszögéből nézve): „Azok voltak a legjobb idők; azok voltak a legrosszabb idők. A bölcsesség kora volt az; meg a bolondság kora. A hit korszaka; a hitetlenség korszaka. A Világosság ideje; a Sötétség ideje. A remény tavasza; a kétségbeesés tele. Előttünk volt még minden; nem volt már előttünk semmi. Mind egyenesen a Mennysországba indultunk; vagy ellenkező irányba.” (Két város regénye, 1859) Az biztos állítható, hogy a mesterséges intelligencia által keltett zűrzavar és ígéretek bizonytalanságot keltenek a jövővel kapcsolatban. A tanulmány szerzője szerint az mesterséges intelligencia rendszerek fejlődése következtében alakot öltő új világot három lehetséges MI-jövőkép testesíti meg – erősen a veszélyekre összpontosítva:

- az MI-rendszerek és az emberek egyesülése, „összeolvadása”, azaz a szingularitás;
- a mindent átható automatizálás utópiája;
- az MI-ágensek által meghatározott világ.

A szingularitást Kurzweil "The Singularity is Near" (2006) című könyve alapján tárgyalja, de ez már az enyhített változat. Újabb, "The Singularity is Nearer" (2024) című könyvében Kurzweil módosítja eredeti elképzelését, a szingularitás nem a szuperintelligens gépek hatalomátvétele lesz, nem fog az emberi faj kipusztulni. Az emberek és a gépek teljes egyesülése fog bekövetkezni, egy új szuperemberi faj veszi át a terepet. Az utópikus jövő-változatban a mindenre kiterjedő automatizálás biztosítja az általános jólétet, végső formájában a kormányzást is automatikus rendszerek végzik, ahogy azt S.M. Sohn leírja „The Last AI of Humanity” című könyvében. Ez a forgatókönyv nem igényli a mesterséges általános intelligencia (AGI) megvalósulását sem, viszont nem vet számot az egyre mélyülő automatizálás által okozott társadalmi, politikai problémákkal. A jövőben ezek egyre nagyobb problémát okoznának, és a végén az emberek alávetett állapotba kerülnének a saját logikájuk alapján működő gépi rendszerekkel szemben. A realista változatban (a szerző a pragmatikus elnevezést használja) olyan MI-alkalmazásokra (MI-ágensek) összpontosítanak, amelyek jobban képesek bizonyos feladatokat elvégezni az embereknél, és ezzel lehetővé teszik, hogy olyan feladatokra összpontosuljon az emberi munkavégzés, amelyeket a gépek nem tudnak jól elvégezni. Jensen Huang, az Nvidia vezérigazgatója szerint ez az MI olyan, mint egy időgép. Néhány másodperc alatt megvalósíthat egy olyan jövőt, amelynek megvalósítása egyébként rengeteg munkát igényelne. A gyors munkavégzés azonban önmagában nem minden. Ezeknek a rendszereknek megbízhatóknak is kell lenniük. A bizalomnak vannak mérnöki megvalósítható aspektusai, de vannak olyan vonatkozásai is, amelyeket a gépi rendszerek (legalábbis jelenleg) nem képesek kielégíteni. A cikk az MI fejlődése kapcsán keletkező problémák terjedelmes katalógusával zárul. Ezek egyelőre nehezen megoldhatónak tűnnek, és ennek következtében Dickens idézete tökéletesen illusztrálja a mai helyzetet.

[Three AI Futures](#)**Kína mesterséges intelligencia jövőképe más – és lehet, hogy okosabb**

Az Egyesült Államokban a digitális ipar törekvéseinek középpontjában a mesterséges általános intelligencia (AGI) megvalósítása áll. A feltételezések szerint az emberi képességeket meghaladó rendszerek megjelenése lenne a következő nagy ugrás az emberi társadalom fejlődésében, ezért aztán dollár milliárdokat fordítanak erre a célra. Kínában ezzel szemben a hozzáállás sokkal gyakorlatiasabb, a mindennapi feladatok tömegében jól használható alkalmazásokra összpontosítanak. A nagy hatású MI-alkalmazásokat nem a jövőbeli elméletek tárgyának tekintik, hanem olyanoknak, amelyet itt és most lehet használni. Természetesen az amerikai technológiai vállalatok is számos gyakorlati alkalmazást fejlesztenek a jelenlegi MI segítségével, de az Egyesült Államokkal ellentétben, ahol az iparág szereplői nagyrészt saját magukra hagyatkoznak, Kínában az állam teljes erejét, eszközeit (nem csak pénzügyi értelemben) mozgósítják a célok eléréséhez – hasonlóan, ahogy az elektromobilitás és a zöld gazdaság esetében történik. „AI+” a neve annak az átfogó kínai stratégiának, amely erőteljesebb összkormányzati erőfeszítéseket irányoz elő a mesterséges intelligencia integrálására a tudományos és technológiai kutatásba, ipari fejlesztésbe és más alkalmazói területek folyamataiba. A cél az, hogy 2030-ra a mesterséges intelligencia átfogóan és mélyen áthassa és támogassa Kína gazdasági fejlődését. Ez nem azt jelenti, hogy egy-egy kínai vállalat nem érdekelt az AGI-hoz kapcsolódó fejlesztésekben, de a kínai rendszer sajátosságaiból következő állami, társadalmi lehetőségek fókuszában a mindennapi alkalmazás áll. Úgy is fogalmazhatunk, hogy hagyják, hogy technológiai vezető szerepe birtokában az Egyesült Államok viselje a felfedezés és a kutatás költségeit, majd Kína az adottságaira támaszkodva gyors követőként optimalizálja a rendszereket a gyakorlati megvalósítás és használat érdekében. Ennek keretében Kína egyre aktívabban alkalmazza a nyílt forráskódú modelleket, amelyek ingyenesen letölthetők és módosíthatók a felhasználók számára, így olcsóbb és könnyebb a kínai vállalatok számára, hogy e technológia köré építsék vállalkozásaikat. Ez a megközelítés segíti a kínai mesterséges intelligencia globális elterjedését is, ez a trend pedig a Szilícium-völgyet is a példájuk követésére készíti.

[China Has a Different Vision for AI. It Might Be Smarter](#)

MI adatközpontokat „zöldít” egy akkumulátor újrahasznosító üzem

A környezeti kérdések nem csupán általában jelentik a közélet és a nemzetközi politika egyik súlypontját. Jó ideje erősödik annak a felismerése is, hogy a társadalmakat minden téren átformáló mesterséges intelligencia technológiák egyik nem várt hátulütője a fejlesztésükhöz, tréningezésükhöz, majd működtetésükhöz kapcsolódó számottevő környezeti „lábnyom”, azaz, a természetre nehezedő plusz terhelés. Az MI adatközpontok különösen sok fejfájást okoznak az üzemeltetőknek, kettős értelemben is. Egyrészt nagyon jelentős közvetlen környezetterhelő hatásuk van, elsősorban a jelentős hőkibocsátás révén. Másrészt viszont szinte elképesztő mértékben „hörpölik” működésük során az energiát; nem véletlen, hogy az adatközpont üzemeltetők közül többen foglalkoznak meglevő atomerőművek megvásárlásával, ill. leállított nukleáris erőművek újraindításával. Az Egyesült Államok sivatagos régiójában, a nevadai Reno városában egy érdekes kísérleti üzem keres (legalább részleges) megoldást erre az egyre súlyosabb problémára. A projekt külön szellemessége, hogy egy környezetvédelmi célú üzem működését alakítják úgy, hogy ebben a másik, előbb említett a mesterséges intelligencia-működtetéshez kapcsolódó környezetvédelmi gondokra biztosítson gyógyírt. Az ottani recikláló üzem elektromos járművekből kiselejtezett akkumulátorok újrahasznosítását végzi. Nemrégiben felfedezték, hogy a kidobott, nagy teljesítményű akkumulátorok egy része még töltési kapacitásának mintegy a felét őrzi; szétbontásuknál tehát lényegesen nagyobb hasznót hajthatnak azzal, ha ezt a még meglevő elektromos kapacitást inkább kinyerik és hasznosítják. A telepen összegyűlt, még felhasználható akkumulátorokból ún. mikrohálózatot alakítanak ki, csatlakozva a recikláló üzem tulajdonosa által már eddig kialakított, napelemekkel táplált mikrohálózatokhoz. Ez a kisléptékű energiaellátó rendszer jelenleg egy Crusoe nevű közeli cég számára biztosít energiát. A Crusoe kriptovaluta bányászattal foglalkozik, ám nagyobb fejlesztések kulcsszereplője is: mintegy 500 milliárd dollár értékben épít MI-adatközpontokat az OpenAI cég számára.

[This battery recycling company is now cleaning up AI data centers](#)

Vita a mesterséges általános intelligencia (AGI) örvén

A mesterséges intelligenciáról szóló vitákban a mesterséges általános intelligenciát (AGI) – az emberihez hasonló, vagy azt meghaladó intelligenciát – úgy kezelik, mintha az a technológiai fejlődés folyamatában egy végső, döntő szint lenne. A vita azonban nem annyira tudományos vitára hasonlít, mint inkább egy névhasználati jogokért folyó versenyre, mert az „emberihez hasonló intelligencia” fogalmának nincs természetesen pontosan meghatározható határa, amelyhez viszonyíthatnánk. Ez önmagában nem baj, hiszen sok más fogalom – például bölcsesség, élet, faj – határai is elmosódtak. A pontos definíció hiánya nem jelenti azt, hogy a kérdéses dolog lehetetlen; az, hogy nem tudjuk megkülönböztetni az embert a géptől, nem jelenti azt, hogy az AGI (vagy az emberihez hasonló intelligencia) nem érhető el. Éppen ellenkezőleg, ez a bizonytalanság arra utal, hogy az emberi és a gépi intelligencia egy folytonos skálán helyezkedik el. Az „emberi” / „nem emberi” megkülönböztetése nem egy természetes skálán élesen jelentkező szakadás és ez a skála társadalmilag konstruált. A különböző közösségek különböző dolgokat várnak el az „AGI”-től: logikai gondolkodást, érzelmi tudatosságot, kreativitást, autonómiát. Közös mérce nélkül viszont az AGI-viták könnyen retorikai csetepatékba fulladnak. A nagy nyelvi modellek sok feladatban már szinte emberhez hasonló teljesítményt nyújtanak, sőt bizonyos standardizált vizsgákon még a legtöbb vizsgázó embert is

felülmúlják. Egyesek szerint ezek a modellek már az AGI kategóriájába tartoznak, mások viszont „kifinomult papagájoknak” tartják őket, amelyek valójában még mindig nem értik, amit mondanak. A nézeteltérés nem igazán a tényekről szól, hanem a mércéről. Mivel nincs egyetértés a képességek küszöbértékeiről, a vita olyan, mintha különböző térképek alapján vitatkoznánk a koordinátákról: a megoldás lehetetlen. Helyesebb lenne nem azon vitatkozni, hogy egy adott rendszert minek nevezzünk. A lényeges kérdés, hogy mire képes. Egyértelmű küszöbértékeket kellene meghatározni – amikor egy MI-rendszer elér egy bizonyos szintet a kulcsfontosságú területeken, akkor ennek megfelelő szabályozási és etikai felülvizsgálatot kell kiváltania, függetlenül attól, hogy minek nevezzük az adott rendszert. Az ilyen küszöbértékeknek mérhetőnek kell lenniük a következő jellemzők esetében: a rendszerek döntéshozatali autonómiája, több területet felölelő tanulás képessége, és annak kiszámíthatatlansága, hogy mennyire fogják követni az emberi utasításokat. Ha ezeket az értékeket átlépik, akkor a társadalmi következményekkel kell foglalkoznunk. Turing megjósolta, hogy egy napon teljesen természetes lesz arról beszélni, hogy a gépek gondolkodnak, anélkül, hogy ez ellentmondásra ingerelne valakit is. Talán a nem túl távoli jövőben az „AGI” kifejezés feleslegessé válik. Nem az lesz fontos, hogy az intelligencia melyik kategóriájába tartozik valami, hanem az, hogy mit jelent ez a társadalmaink számára, és hogy van-e elég bölcsességünk ahhoz, hogy együtt éljünk vele.

[AGI: a battle over names, not minds](#)





Működésben

A generatív MI-vel támogatott programozás rejtett költségei

A generatív mesterséges intelligencia robbanásszerűen terjed a tudásalapú munkákban, különösen a szoftverfejlesztésben. Az OpenAI GPT-4.1 verziója nagy hangsúlyt fektet a kódolási képességek fejlesztésére, és egy lépéssel közelebb visz a teljes automatizáláshoz. Azok a szervezetek, amelyek bevezetik ezeket az eszközöket, jelentős nyereségre számítanak. A GitHub jelentése szerint a Copilotot használó programozók akár 55%-kal is termelékenyebbek, a McKinsey pedig megállapította, hogy a fejlesztők a generatív mesterséges intelligencia segítségével akár kétszer olyan gyorsan tudják elvégezni a feladataikat. De ezeknek a pozitív mutatóknak van egy fontos korlátozó tényezője. A tanulmányokat ellenőrzött környezetben végezték, ahol a programozók izolált feladatokat hajtottak végre – nem pedig a való világban, ahol a szoftvereket komplex, meglévő rendszerekre kell építeni. Az MIT Sloan Management Review közölt érdekes tanulmányt ebben a témában, ahol a szerzői gárda egyértelműen leszögezi: a generatív mesterséges intelligencia nagyon hatékony teljesítmény növelő eszköz lehet, de csak akkor, ha kellő körültekintéssel alkalmazzák. Ha nem a szükséges gondossággal történik az MI programozó képességeinek felhasználása, akkor nagyon könnyen a tényleges méretfejleszthetőség (skalázhatóság) láthatja ennek a kárát. Arról nem is beszélve, hogy az ilyen (átgondolatlan, vagy éppen kódhibákat tartalmazó) szoftverek vagy szoftver elemek tökéletesen destabilizálhatják azokat a rendszereket, amelyekre ráépülnek. A szerzők kutatásuk során interjúkat készítettek szoftverfejlesztésben részt vevő személyekkel – kezdő fejlesztőktől vezető szoftvermérnökökig és felsővezetőkig – különböző iparágakban, többek között a biztosítási, webtárhely-szolgáltatási, közösségi média, védelmi, menedzsment tanácsadási és fintech szektorokban. Ezekből az interjúkból, a szaklapok áttekintéséből és saját gazdasági modellezésükből kiindulva azonosítottak több olyan stratégiai szempontot, amelyet a vállalatoknak figyelembe kell venniük, amikor generatív MI-t alkalmaznak a szoftverfejlesztéshez.

[The Hidden Costs of Coding with Generative AI](#)

Algoritmikus segédedző dolgozik az Országos Amerikaifutball Liga (NFL) trénerének keze alá

A sport társadalmi jelentősége mindenütt nagy, és ez óhatatlanul előrevetíti ennek a területnek a fontosságát az MI-technológiák adaptálásában, legyen szó szélesan értelmezett elfogadásról, vagy szűkebben értelmezett üzleti sikerről. Az Egyesült Államokban a nálunk sokáig rögbinek nevezett sportág a „sportok sportja” – persze a szintén rendkívül népszerű baseball mellett. A sport országos szintű bajnoki mérkőzés sorozata, az NFL (National Football League) tízmilliókat tart lázban Amerikaszerzte; ennek megfelelően hallatlan összegeket megmozgató üzleti ökoszisztéma épült köré, de a társadalmi-belpolitikai jelentősége önmagában is szinte páratlan. Az MI-fejlesztések híreit figyelőknek nem újdonság, hogy a sport, különösen az élsport (vagy más megfogalmazással a sport biznisz) hosszú ideje az egyik fontos alkalmazási terepe a különféle MI-technológiáknak. A sport ugyanis – a leendő sportolók kiválasztásától a tréningezésen át a mérkőzések sikerességének garantálásáig – bizonyos értelemben tudományos tevékenységek sora, összessége. Sok évtizede láthatjuk, hogy a mindenkori technológia (kezdve a filmfelvevőktől, a folyamatok visszajátszására, utólagos elemzésére alkalmas videó eszközökön át a számítástechnika egyre újabb és újabb vívmányaiig) hogyan épül be a sport, a sportüzlet mindennapjaiba. A pusztá emberi teljesítmény, önmagában, már jó ideje nem elég a modern sportban való kiemelkedésre, csúcsteljesítményekre. Az NFL kiemelkedő csapatai, elsősorban a sportolók felkészítéséért felelős edzők 2014 óta használják munkájuk során a Microsoft cég táblagépeit. A rajtuk futtatott programok a csapatok felkészítésén át a meccsek hatékony menedzseléséig sok területen igyekeztek segíteni a sportembereket – a hírek szerint sokszor nagyon vegyes eredményességgel. Az első mobil számítástechnikai eszközök mellett éppen ezért sohasem tűnt el teljesen a szakvezetők, edzők kezéből a papír és a ceruza, és a „digitalizáció” is jobbára az Excel-alapú információ kezelésben csúcsonosodott ki. A napokban az NFL vezetése újabb szerződést kötött a Microsofttal: ennek keretében az első osztály kiemelkedő csapatainak vezérkara Qualcomm-alapú Copilot+ táblagépeket kap használatra. A mesterséges intelligencia-alapú edzői munkában az MI-segédedző a legkülönbélebb helyzetekben hivatott gyors, szakszerű támogatást nyújtani a felkészülésnél vagy a játékban.

[NFL coaches' new AI assist](#)

MI-alapú megfigyelő rendszer védelmezi az Indonézia körüli veszélyes tengerek biztonságát

Az Indonéz Parti Őrség a közelmúltban 10 évre szóló együttműködési megállapodást kötött egy angliai technológiai céggel, az SRT Marine Systemmel annak érdekében, hogy a legkorszerűbb szenzorokat hálózatba kapcsoló, és a monitorozó tevékenységet mesterséges intelligencia szoftverrel támogató tengeri biztonsági rendszert építsenek ki, és működtessenek. A 158 millió euró értékűre becsült szerződéssel Indonézia egy olyan hatékony tenger-felügyeleti rendszerre szeretne szert tenni, amely nem csupán a roppant veszélyes kalózkodást igyekszik majd felderíteni, nyomon követni, megelőzni. A rendszer kiemelt feladatai közé számítják ugyanis a környező tengerek ökológiai biztonságának védelmét is. A tervek szerint a műholdas, illetve földi, tengeri érzékelőkből származó adatok feldolgozásával a mind súlyosabb környezetvédelmi fenyegetések ellen is hatékonyan tud majd fellépni Indonézia.

[New AI surveillance system to monitor Indonesia's Seas](#)

A ChatGPT legutóbbi modellváltásával eltűnő régi verziót gyászoló felhasználók

A nagy nyelvi modelleket használó szakmai és amatőr közönség széles körében váltott ki számottevő mértékben reakciókat a területen talán mindmáig piacvezető OpenAI cég legutóbbi modellfrissítő akciója. A szokás szerint hosszú ideje sejtetett, és éppen ezért nagy izgalommal várt ChatGPT-5 verzió azonban legfeljebb félig sikerült újításnak nevezhető – és az is talán csak jóindulattal. A felhasználók tömegei tapasztalták úgy, hogy az 5-ös modell nem jobb, erősebb, vagy nagyobb képességű lett elődeinél, hanem egyenesen visszalépést hozott azokhoz képest. Számunkra itt nem is a vita maga az érdekes, hanem a modellfrissítés, illetve az új modellel kapcsolatos felhasználói percepciók nyomán kibontakozó érzelmi reakciók. A téma jelentőségét, társadalmi fontosságát érzékeltetendő gyorsan leszögezzük: az „érzelmi reakciók” alatt itt nem bosszúságot, csalódottságot, elégedetlenséget, mérget, esetleg hasonló jellegű, de pozitív megnyilvánulást kell érteni. A mértékadó tudományos műhely, az MIT technológiai-fejlesztésekkel foglalkozó központjának elemzése azt mutatja, hogy – a felfokozott várakozásaikban csalódott – felhasználók nem elhanyagolható rétegében szabályos gyászt váltott ki az OpenAI akciója. A frissítéssel, tehát az 5-ös modell közzétételével egy időben az OpenAI elérhetetlenné tette ugyanis a ChatGPT korábbi, sokféle és nagyon népszerű változatait, mint például a ChatGPT-4o modellt. Felhasználók egy része élte meg ezt a váltást úgy, mintha egy közeli emberi kapcsolatát, rokont, vagy barátot veszített volna el. Szórakoztató anekdotánál itt messze többről van szó. Ez a különös fejlemény (is) egyértelműen mutatja, hogy az új MI-technológiák, elsősorban természetesen az LLM-ekre épített chatbotok olyan erős, korábban kizárólag ember-ember közötti viszonyban kiépülő kapcsolódást képesek előidézni, ami merőben új jelleget adhat a most épülő ember-gép konstellációnak. Arról van tehát szó, hogy az új nagy nyelvi modellek, és más generatív technológiák fejlesztőinek, üzembe állítóinak nagyon komolyan tekintetbe kell most már venniük azt, hogy az új technológia érzelmi síkon képes újfajta viszonyrendszert létrehozni a felhasználókkal. Ebben az értelemben az MI – a fejlesztők akaratától szinte függetlenül – messze több, mint egy eszköz: a társadalmi szövetnek egy darabja maga is.

[Why GPT-4o's sudden shutdown left people grieving](#)

