



JÖVŐ

/ gépek • képek

avagy jövőnk az AI korában

KONFERENCIAKÖTET

Tudományos és üzleti horizontok konferenciája

INTERDISZCIPLINÁRIS KONFERENCIA

Veszprém • VEAB Székház
2025. december 2.

MTA-VEAB Gazdaság-, Jog- és Társadalomtudomány
Szakbizottság és az MI és a Kreatív Iparágak
Társadalomtudományi Kutatási Munkabizottság

JÖVŐKÉPEK- JÖVŐGÉPEK, AVAGY A JÖVŐNK AZ MI FÉNYÉBEN

Konferenciakötet

2025. december, Veszprém

VEAB; MI és Jövő kutatások munkabizottsága

Felelős Kiadó: Molnárné Dr. László Andrea

ISBN 978-615-02-7338-9

2026.

Tartalomjegyzék

Tartalom

A mesterséges intelligencia hatása a gazdasági növekedésre és a termelékenységre.....	1
---	---

Bögel György 		1
Absztrakt.....		1
Abstract.....		1
1. Rövid áttekintés az általános célú technológiákról.....		2
2. A dotcom válság és következményei.....		5
3. A GDP és a termelékenység növekedése a hatvanas évek óta.....		6
4. A „termelékenységi paradoxon”.....		8
5. Fellendülés a dotcom válság előtt és után.....		10
6. Korunk általános célú technológiája: a mesterséges intelligencia.....		12
7. Következtetések és kérdőjelek.....		15
Irodalomjegyzék.....		18
A tudatos mesterséges intelligencia kora: ember és gép közös döntéseinek szabályozási és gazdasági kérdései.....		20

Halász Rita 		20
Absztrakt.....		20
Abstract.....		20
1. Bevezetés.....		21
2. A mesterséges intelligencia gazdasági és társadalmi jelentősége.....		21
3. Az ember és a mesterséges intelligencia együttműködésének fő paradigmája.....		22
4. A mesterséges intelligencia gazdasági szerepe.....		23
5. A vállalati AI-alkalmazások fejlődési szintjei.....		24
6. A mesterséges intelligencia alkalmazásának kockázatai.....		25
7. Nemzetközi AI-tanulmányok és kutatási eredmények: gazdasági hatások, szabályozási trendek és az ember–AI együttműködés.....		27

8. A mesterséges intelligencia szabályozásának normatív alapjai: a megbízható AI koncepciója.....	28
9. Az Európai Unió és Magyarország mesterséges intelligencia szabályozása, és a kockázatalapú modell gyakorlati érvényesülése	30
10. Az Egyesült Államok és Kína AI-szabályozásának eltérő megközelítései.....	32
11. Az ember szerepe a mesterséges intelligencia korszakában.....	33
12. Következtetések	35
Irodalomjegyzék.....	36
Algoritmusok a csatatéren: A mesterséges intelligencia architektúráinak hatása a modern hadijogra és a döntéshozatali mechanizmusokra	38
Vécsey Richárd Ádám 	38
Absztrakt.....	38
Abstract.....	38
1. Bevezetés.....	39
2. A mesterséges intelligencia technológiai és matematikai alapjai	39
2.1. A gépi tanulás matematikai alapjai.....	39
2.2. Mélytanuló architektúrák katonai alkalmazásokhoz.....	40
2.3. Adat- és szenzorarchitektúrák.....	41
3. A katonai döntéshozatal átalakulása: OODA és OODA(B)	42
4. Etikai és stratégiai kérdések.....	43
4.1. A belépési küszöb csökkenése.....	44
4.2. Túlzott bizalom és automatizációs torzítás	44
4.3. A háború köde és a dehumanizáció.....	44
4.4. Hiperháború és a műveleti tempó gyorsulása	45
5. A katonai mesterséges intelligencia és a hadijog	45
5.1. Jus contra bellum	45
5.2. Jus ad bellum.....	46
5.3. Jus in bello.....	46
5.4. Jus post bellum.....	46

6. Összegzés, a jövő irányai	47
6.1. Predikció: a jövő előrejelzése mint katonai képesség	47
6.2. Prevenció: a konfliktusok megelőzése és a korai beavatkozás dilemmái	47
6.3. Eszkaláció: a gyorsuló hadviselés kockázatai	47
6.4. Katalizáció: az MI mint erőszorzó és rendszerszintű átalakító tényező	48
6.5. Redempció: felelősség, helyreállítás és a háború utáni igazságosság kérdése	48
Irodalomjegyzék	48
Mesterséges intelligencia a levegőben	50
Katona Róbert 	50
Absztrakt	50
Abstract	50
1. Bevezetés	51
2. MI fogalma	52
3. Mesterséges intelligencián alapuló dróntechnológiákkal kapcsolatos elvárások	52
4. Az autonóm rendszerekkel szembeni elvárások	53
5. Az autonóm technológiák szervezeti hatásai és a tanácsadói szektor szerepe	54
6. Dróntechnológia előnyei a hagyományos módszerekkel szemben	55
7. Biztonságos leszállás az MI alapú VTOL segítségével	55
8. Összefoglaló	56
Irodalomjegyzék	56
Felépíteni az ismeretlent – Autonóm, kooperatív jövőgépek és digitális ikrek a fenntartható földön kívüli infrastruktúra-fejlesztésben	58
Building the Unknown – Autonomous, Cooperative Future Machines and Digital Twins in Sustainable Extraterrestrial Infrastructure Development	58
Dr. Ország-Krisz Axel 	58
Absztrakt	58
Abstract	58
1. Bevezetés	59

1.1. Az in-situ erőforrás-hasznosítás (ISRU) gazdaságtana és logisztikája	59
1.2. A Föld-centrikus ellátási láncok fizikai korlátai	59
2. A Földön Kívüli Környezet Determinisztikus Modellezésének Korlátai.....	60
2.1. Regolit-dinamika: koptató hatás, elektrosztatikus feltöltődés és vákuum-kohézió	60
2.2. A klasszikus hibakorrekció és a humán intervenció hiányának rendszerszintű kockázatai	60
3. Elméleti Módszertan: Digitális Iker és Kockázatkezelő Mesterséges Intelligencia	61
3.1. A Digitális Iker architektúrája és valós idejű szinkronizációja	61
3.2. Curriculum Learning (Tanterv-alapú tanulás) a Markov-féle döntési folyamatokban (MDP)	62
3.3. Sztochasztikus MI: Episztemikus és aleatórikus bizonytalanságok számszerűsítése (Bayesi neurális hálók)	62
4. A Földön Kívüli Infrastruktúra-Építés 7 Lépéses Folyamatmodellje.....	63
1. Lépés: Ontológiai és tudományos tudástár felépítése.....	64
2. Lépés: Digitális iker generálása és szimulációs környezet építése	65
3. Lépés: A bizonytalanságra tanított és biztosított MI integrációja.....	65
4. Lépés: Kooperatív robotikai protokollok konfigurálása	65
5. Lépés: Hitelesítés, szabványosítás és ellenőrzés.....	65
6. Lépés: Kétlépcsős tesztelés és validáció	65
7. Lépés: "Go - No Go" döntéshozatal és autonóm kivitelezés	66
5. Esettanulmány: Holdbéli Menedék- és Leszlóhely Létesítése	66
5.1. A probléma felbontása (Work Breakdown Structure - WBS) és geometriai kritériumok	66
5.2. Heterogén flotta-architektúra: Felderítők, nehézgépek és szuborbitális drónok kooperációja	66
6. Globális Technológiai Körkép és TRL-Szintek Kritikai Elemzése.....	67
6.1. NASA RASSOR, DLR ARCHES és az Etna-analóg missziók tanulságai	67
6.2. Additív gyártástechnológiák regolit-alapon: Regolight és 3D Habitat Challenge	68
6.3. Additív technológiák és strukturális konszolidáció: Regolight és NASA Habitat Challenge	68

7. Összegzés.....	68
Felhasznált irodalom.....	69
A generatív mesterséges intelligencia hatása a szervezeti kultúrára: kommunikáció, tanulás és vezetői felelősség.....	71
The Impact of Generative Artificial Intelligence on Organizational Culture: Communication, Learning and Leadership Responsibility.....	71
Molnárné László Andrea 	71
Absztrakt.....	71
Abstract.....	71
1. Bevezetés.....	72
2. Szakirodalmi háttér.....	72
2.1. Az MI mint szociotechnikai és kulturális változás.....	72
2.2. MI-közvetített kommunikáció, hitelesség és nyelvi normák.....	73
2.3. Munkahelyi tanulás és a kompetenciák átrendeződése.....	73
2.4. Algoritmikus menedzsment és vezetői felelősség.....	74
3. Módszer.....	75
4. Eredmények és értelmezés.....	75
4.1. Az emberközpontú értékek újrafókuszálása.....	75
4.2. A kommunikáció hatékonysága és a hitelesség közötti feszültség.....	75
4.3. A tanulásközpontú kultúra lehetősége és feltételei.....	76
4.4. Vezetői kontroll, jóllét és felelősség.....	76
4.5. Innováció: gyorsítás, kísérletezés és homogenizálódás.....	77
5. Vezetői következtetések és alkalmazási keret.....	77
6. Összefoglaló.....	78
Irodalomjegyzék.....	79
Nagy Nyelvi Modell (LLM) alapú szoftver az időskor szolgálatában – A hosszú és minőségi élet szoftveres támogatása.....	81
Large Language Model (LLM)-based software for the elderly – Software support for a long and quality life.....	81

Dr. Ország-Krisz Axel 	81
Absztrakt.....	81
Abstract.....	82
1. Bevezetés és Irodalmi Áttekintés.....	83
1.1. Demográfiai kontextus: A társadalom előregedése és az ellátórendszerek krízise	83
1.2. A gerontechnológia evolúciója: A passzív segélyhívóktól a generatív mesterséges intelligenciáig.....	83
1.3. Proaktív vs. reaktív gondoskodás: A preventív intervenció szükségessége.....	84
1.4. Jelen tanulmány célkitűzése: A javasolt rendszer elméleti keretei.....	84
2. A valós idejű beszédalapú interakció architektúrája és szűk keresztmetszetei.....	85
2.1. A társalgási késleltetés (Latency) anatómiája.....	85
2.2. Időskori automatikus beszédfelismerés (ASR) specifikus akadályai	86
2.3. Szövegfelolvasás (TTS) és szinkronizáció	87
3. Kognitív állapotfelmérés és mentális tesztek az LLM-ben: Kompetencia és felelősség.....	88
3.1. Strukturálatlan beszélgetésből nyert kognitív biomarkerek.....	88
3.2. Validált klinikai tesztek implicit integrációja.....	88
3.3. Jogi és bioetikai felelősségi körök	89
4. A fizikai aktivitás multimodális értékelése és validációja	90
4.1. A fizikai státusz monitorozásának technológiai korlátai.....	90
4.2. Az LLM mint kontextuális validációs motor	91
4.3. Személyre szabott, adaptív mozgástámogatás	92
5. Kitekintés: Az asszisztens mint a hosszú egészséges élettartam (Longevity) eszköze	93
5.1. A Lifespan (élettartam) és a Healthspan (egészséges élettartam) szétválása.....	93
5.2. Epigenetikai és életmódbeli optimalizáció éjjel-nappal	93
5.3. A szociális izoláció mint biológiai öregedési faktor	94
6. Összegzés és jövőbeli kutatási irányok	95
6.1. A tanulmányban bemutatott elméleti és mérnöki keretek szintézise	95
6.2. Jövőbeli kutatási irányok és technológiai mérföldkövek	96

6.3. Záró gondolatok.....	96
Felhasznált irodalom	97

A mesterséges intelligencia hatása a gazdasági növekedésre és a termelékenységre

Bógel György¹ 

Absztrakt

A mesterséges intelligencia századunk elejének valószínűleg legfontosabb általános célú technológiája. A korábbi nagy innovációs hullámok lefutása a lappangástól kezdődően a beérésig jellegzetes mintázatot mutat. Feltételezhetjük, hogy a modern mesterséges intelligencia esetében is hasonló fázisokat és jelenségeket figyelhetünk meg. A tanulmány első fejezete áttekintést ad az általános célú technológiák által gerjesztett innovációs hullámok sajátosságairól. A második szakasz összehasonlítási céllal részletesen elemzi az ezredfordulón jelentkező internetes hullámot és annak következményeit. A harmadik bemutatja, hogyan nőtt a világgazdaságban a GDP és a termelékenység a hatvanas évektől napjainkig. A negyedik összefoglalja az úgynevezett „termelékenységi paradoxon” lényegét. Az ötödik fejezet a dotcom válságot követő fellendülés okaira világít rá, felhívva a figyelmet a késleltetett hatások jelenségére. A hatodik a modern mesterséges intelligencia fejlődését és elterjedését párhuzamba állítja a gazdasági növekedéssel. A tanulmány legfontosabb következtetése az, hogy pozitív jelek ellenére a makrogazdasági mutatókban radikális javulást egyelőre nem tapasztalhatunk, az MI hatása még nem bontakozott ki, és ahhoz, hogy ez megtörténjen, számos feltételt kell megteremteni, akárcsak az internetes hullám idején – erről szól a 7. fejezet. Mivel a technológiai fejlődés rendkívül gyors és a világgazdaságban is váratlan események jöhetnek, jeleznünk kell, hogy a kéziratot 2026 februárjának elején zártuk le. A tanulmány az MTA-VEAB és a Mind Mate Inspiration által közösen rendezett, 2025. december 2-án Veszprémben tartott konferencián elhangzott hasonló témájú előadás bővített és friss adatokkal dúsított változata. Szerzője köszönettel tartozik a rendezvény szervezőinek és a kapcsolódó tanulmánykötet szerkesztőjének áldozatos munkájukért.

Kulcsszavak: általános célú technológiák, mesterséges intelligencia, innováció, gazdasági növekedés, termelékenység

Abstract

Artificial intelligence is probably the most important general-purpose technology of the beginning of this century. The course of previous major waves of innovation from latency to maturity shows a characteristic pattern. We can assume that similar phases and phenomena can be observed in the case of modern artificial intelligence. The first chapter of the study provides an overview of the characteristics of innovation waves triggered by general-purpose technologies. The second section analyzes in detail the Internet wave that emerged at the turn of the millennium and its consequences for comparative purposes. The third shows how GDP and productivity have grown

¹ Research Affiliate, Central European University. Elérhetőség: bogelgy@ceu.edu.

in the world economy from the sixties to the present day. The fourth summarizes the essence of the so-called “productivity paradox”. The fifth chapter sheds light on the reasons for the recovery after the dot-com crisis, drawing attention to the phenomenon of delayed effects. The sixth parallels the development and spread of modern artificial intelligence with economic growth. The most important conclusion of the study is that, despite positive signs, we cannot yet experience a radical improvement in macroeconomic indicators, the impact of AI has not yet unfolded, and for this to happen, a number of conditions need to be created, just like during the Internet wave – this is what Chapter 7 is about. Since technological development is extremely fast and unexpected events can occur in the global economy, we must indicate that the manuscript was closed at the beginning of February 2026. The study is an expanded and updated version of a presentation on a similar topic given at the conference held in Veszprém on December 2, 2025, jointly organized by the Hungarian Academy of Sciences-VEAB and Mind Mate Inspiration. The author is grateful to the organizers of the event and the editor of the related study volume for their dedicated work.

Keywords: General purpose technologies, artificial intelligence, innovation, economic growth, productivity

1. Rövid áttekintés az általános célú technológiákról

Az általános célú technológiák a gazdaság szinte minden szektorában megjelennek, de a hatásuk számos más területre is kiterjed, így például a társadalmi élet jelenségeire, a társadalom szerkezetére, a fogyasztási szokásokra, a politikai életre, a kultúrára és az oktatásra. Példát keresve elegendő az elektromosságra mint általános célú technológiára gondolni: a modern világban az elektromos eszközök szinte mindenütt jelen vannak a gyáraktól kezdve az otthonokon át a kórházakig.

A széles körű elterjedés mellett az ilyen technológiák másik fontos tulajdonsága a folyamatos fejlődés. Az idő előrehaladtával a technológia, illetve annak egyes elemei egyre olcsóbbak és hatékonyabbak lesznek. Extrém példának tekinthetők a modern számítástechnikában és távközlésben nélkülözhetetlen tranzisztorok, amelyek mérete és egységköltsége Moore törvényét² követve az eredeti apró töredékére zsugorodott.

A harmadik fontos tulajdonság (vagy mondjuk inkább képességnek) az innováció felszabadítása. Ahol az ilyen technológiák felbukkannak, friss impulzusokat kap az innováció, kapuk nyílnak meg az alkotói fantázia előtt, új mérnöki terek bontakoznak ki, új termékek és szolgáltatások kerülnek ki a piacra. Az internet megjelenése és terjedése például lehetővé tette az online kereskedelem és a közösségi média fejlődését.

Széles körű elterjedés, folyamatos fejlődés, az innováció felszabadítása – a történelem során számos ilyen tulajdonságokkal és képességekkel bíró általános célú technológia bontakozott ki és nyomta rá a bélyegét különböző korszakokra. Az 1. sz. táblázat közülük sorolja fel a fontosabbakat a 18. század végétől elindulva napjainkig bezárólag. Kétségtelenül sok benne a bizonytalanság és a kérdőjel. A felsorolásból önkényesen kihagytuk például a genetikát³, a digitális infokommunikációs technológia történetét pedig három szakaszra bontottuk, jelesen a számítógépek, az internet és a mesterséges intelligencia világára. Önkényesek a korszakhatárok is. A belső égésű motoroknál például a 19. század végét jelöltük meg kezdetként, pedig a prototípusok már a század legelején megjelentek.

Egyes általános célú technológiák idővel kifulladásra, eltűnnek (gőzgépeket ma valószínűleg nagyon kevesen használnak), mások viszont sokáig velünk maradnak. Az internet „hőskorának” a

² Moore, 1965

³ Karikó, 2023; Suleyman, 2023

20. század vége és a 21. eleje tekinthető, de persze ezt követően sem tűnt el, csak a jelenléte már nem kelt feltűnést, teljesen hozzászoktunk, inkább a hiánya okoz néha problémát.

Egy-egy általános célú technológia kibontakozásához és elterjedéséhez sokféle forrásra (inputra) van szükség. A táblázat harmadik oszlopa ezek közül csak a fontosabbakat sorolja fel, de ennyiből is látható, hogy igazából nem technológiákról, hanem inkább *technológiai nyálábokról*, különböző technológiák együtteséről van szó, amelyek idővel beérik és kiegészítik egymást, összekapcsolódnak egymással. A robbanómotorok például acélból készülnek, benzinnel működnek, az autók pedig gumibroncsokon gurulnak – ezek előállításához (az egyéb alkatrészekről és tartozékokról nem is beszélve) nagyon különböző technológiákat kíván. A nagy innovációs hullámok akkor indulnak el, amikor a kulcsfontosságú technológiák komponensei tömegessé, olcsóvá, könnyen hozzáférhetővé, szabványossá válnak.

Az általános célú technológiák által elindított innovációs hullámoknak jellegzetes szabályosságai, mintázatai vannak⁴. Berobbanásukat hosszabb-rövidebb ideig tartó lappangás előzi meg, ami a kiegészítő jellegű technológiák felzárkózásához is szükséges. Az internet fejlődése például a hatvanas években indult el katonai projektként (ARPANET), az első host-to-host üzenetet 1969-ben küldték rajta, a TCP/IP hálózati protokollt 1983-ban fogadták el formálisan, Tim Berners-Lee újításai 1990 körül tették a nagyközönség számára hozzáférhetővé a Világhálót, de az internet gyors térhódítása, használatának tömegessé válása csak a kilencvenes évek közepén következett be.

Mivel az általános célú technológiák gyorsan és sokfelé terjednek, sokak fantáziáját megmozgatják, berobbanásukat követően a piacokon és a tőzsdéken „technológiai lufik” fuvódhatnak fel, amelyek jellegzetes tüneteket mutatnak⁵. Rengeteg új vállalkozás születik, a beruházások és a vállalatok értékelése elszakad a realitástól, a tőzsdei árfolyamok az egekbe szöknek, a befektetőket illúziók és mániák vezérlik⁶, az infrastrukturális beruházások aránytalanul, a keresletet hagyva növekednek, a piacon az eladók diktálnak. A sajtóban és a könyvesboltokban hosszú, akár több évtizedes fellendülést ígérő elemzések, prognózisok kerülnek a nagyközönség elé.

1. táblázat. Általános célú technológiák a 18. század végétől napjainkig

⁴ Perez, 2002; Bőgel, 2002

⁵ Komáromi, 2006

⁶ Kindleberger, 2000

Általános célú technológia	Időszak	Szükséges források / inputok
Gőzgép	18. század vége, 19. század	szén, vas, víz, mérnöki tudás
Elektromosság	19. század vége, 20. eleje	fosszilis energiaforrások, vízenergia, rézkábelek, elektromos hálózat
Belső égésű motor	19. század vége, 20. század	benzin, gumí, acél
Számítógép	20. század második fele	szilícium, tranzisztor, bináris logika
Internet	20. század vége, 21. eleje	üvegszálak kábelek, hálózati protokollok
Mesterséges intelligencia	21. század	nagy adattömegek, adatközpontok, számítástechnikai felhő, neurális hálók

Az ilyen piaci és pénzügyi lufik előbb-utóbb kipukkannak, a kedélyek megnyugodnak, az árfolyamok esnek, a kockázati tőke visszahúzódik, a piac józanabb lesz. Mindeközben sok vállalkozás tönkremegy, elveszti az értékét, akár általános gazdasági recesszió is bekövetkezhet. Ilyen válságra számos példa akad a történelemben. Közülük az egyik legtanulságosabb a 19. század közepén kibontakozott brit vasútépítési mánia: a lufi kipukkant, a Bank of England emelni kezdte a kamatlábakat, sok család elvesztette minden megtakarítását, vállalati bukások és felvásárlások következtek, a tarka vállalkozói világot négy regionális társaság monopóliuma váltotta fel. A bajok az egész brit gazdaságra, sőt később Európa más országaira is áterjedtek, általános gazdasági recesszió következett, munkások ezrei kerültek utcára.

A válságos szakaszok után a vásárlók visszafogottabban, óvatosabban kezdenek viselkedni, a piacon egyre inkább ők diktálnak. A termékek beérnek, szabványosodnak, praktikusabbak lesznek. Már nem a *beszerzésük* a legfontosabb, hanem az ésszerű *használatuk*: a beruházásoknak, fejlesztési programoknak meg kell térülniük. A korábbi csábító jóslatokból, nagy ígéretekben sok minden megvalósul, de az ígéretknél lassabban, fokozatosabban. Az innovátorok alkalmazkodnak az őket körülvevő világhoz, a világ pedig alkalmazkodik hozzájuk.

A kölcsönös alkalmazkodás egyre mélyebb és kiterjedtebb lesz, a gazdaság és a társadalom átalakul, összehangolódik az új technológiákkal, ennek minden pozitív és negatív, szándékolt és nem szándékolt következményével. A technikai újdonságok megszokottá, az élet nélkülözhetetlen feltételeivé válnak. A kiépült új infrastruktúra lendületet ad egyes beérett, lelassult iparágaknak, modernizálási akciók indulnak mindenfelé. A bevált, eredményes gyakorlatok, megoldások tananyagokká, rutinokká kristályosodnak, nő a termelékenység, ami jó hatással van a makrogazdasági eredménymutatókra. A technológiai szektor ilyenkor általában konszolidálódik, a piacot néhány nagyvállalat uralja.

Történelmi tapasztalatok szerint a kijózanodás és összerendeződés szakaszát a beérés fázisa követi. A növekedés lassul, megmutatkoznak a fejlődés korlátai. A pénz, a kockázati tőke új leszállópályát keres, a világ a következő innovációs hullámra vár.

Mivel érdekes tanulságokkal szolgálhat a mesterséges intelligenciára nézve, az *1. sz. táblázatban* felsorolt általános célú technológiák közül válasszuk ki most az internetes hullámot, és nézzük

meg, hogyan jelent meg benne a „technológia lufi”, hogyan szökött fel majd csillapodott le a piaci láz!

2. A dotcom válság és következményei

Mivel ez a történet a közelmúltban játszódtott le óriási médianyilvánosság közepette, a vizsgálódáshoz számtalan adat és beszámoló áll a rendelkezésünkre⁷.

Az internetes lufi megjelenése és felfújódása jól követhető az amerikai technológiai tőzsde, a Nasdaq indexének alakulásán. A múlt század kilencvenes éveinek elején az index 450 pont körül ingadozott. A nyugalomnak 1995-ben szakadt vége: ekkor jelent meg a tőzsdén az előző korszakban nagyra nőtt Microsoft kihívójaként a fiatal Netscape cég, az internetet gyakorlatilag bárki számára könnyen hozzáférhetővé tevő Navigator böngésző fejlesztője. A Nasdaq index a következő öt évben az ötszörösére nőtt: ezerről ötezer pontra. Az általános gazdasági környezet egyébként is kedvező volt: a GDP nőtt, az alacsony kamatlábaknak és a gazdaságban felhalmozódott sok pénznek köszönhetően viszonylag könnyű volt tőkéhez jutni, az internet új innovációs teret nyitott meg sok induló vállalkozás előtt, és ahogy ez már lenni szokott, mindenfelé optimista jóslatok láttak napvilágot töretlen fejlődésről, az előttünk álló, technológiai alapú aranykorról⁸.

A befektetői világ is felélénkült. A beruházásra fordítható pénzek, és persze a kockázati tőke elsősorban a technológiai, és azon belül az internetes szektort (az úgynevezett „dotcom-cégeket”) vették célba, 1998-ban már gyakorlatilag minden második dollár itt talált leszállóhelyet. Tömegével jelentek meg új kockázati alapok, dollármilliárdokat pumpáltak internetes startupokba. A polgárok sem maradtak tétlenek: az amerikai háztartások vagyonában a részvények aránya a kilencvenes évek végére közel megduplázódott⁹.

Az induló internetes vállalkozások általában a növekedésre helyezték a hangsúlyt, a nyereséggel kevésbé törődtek. Irányíts magadra minél nagyobb figyelmet, igyekezz minél gyorsabban minél több ügyfelet szerezni, mindegy, hogy milyen áron, a bevétel majd jön később magától! – ez volt a „stratégia”. Rengeteg pénz égett el a rózsás jövő reményében. A befektetők és a vállalkozók viselkedésébe *irracionális* elemek vegyültek. A befektetéseknél, illetve a tőzsdén sok vállalt ért el milliárdos értékeléseket arányos bevétel és nyereség nélkül (számos cégnek bevétele sem volt, nemhogy nyeresége); az értékelések elszakadtak a hagyományos pénzügyi mutatóktól, a fundamentumoktól.

1998 vége és a rákövetkező év eleje között a lázban égő befektetők 400%-kal nyomták fel az internetes részvények értékét, messze túlszárnyalva a Dow Jones Industrial Average emelkedését. A Nasdaq 2000 márciusában tetőzött 5.048 ponttal. Ez után drámai zuhanás következett, szinte egyik óráról a másikra. A tőzsdén eladási hullám indult el, sok internetes vállalatról kiderült, hogy nem életképes, a király meztelen. A befektetések értéke lezuhant, a Nasdaq index ezer pont körül fékezett le, ami után lassan ismét nőni kezdett, de a korábbi rekordot csak 2015 körül érte el.

A technológiai piac alaposan megváltozott: a vállalati IT beruházásokat visszavágták, sok technológiai vállalat egyszerűen eltűnt a piacról, a megmaradók pedig súlyos veszteségeket voltak kénytelenek elkönyvelni. A befektetők elhidegültek a friss internetes cégektől, a kockázati tőke felszívódott, másfelé fordult. Az egyetemi informatikai szakokra kevesebben jelentkeztek. A drámai események, a nagy csalódások hatására sajátos csalási hullám is elindult, a sajtó

⁷ Lásd pl. Mandel, 1999

⁸ Lásd pl. Schwartz, Leyden, Hyatt, 1999

⁹ Perkins; Perkins, 2001

világraszóló botrányokról (lásd pl. az Enron vagy a WorldCom esetét), börtönbe került vezetőkről tudósított.

A dotcom válság az innováció területén lemaradó Európát csak kevéssé érintette, Magyarországot pedig csak meglegyintette a szele.

A krízis súlyos volt, de nem tartott sokáig. A digitalizálás általános trendje nem állt meg, a számítógépek idővel szinte mindenütt megjelentek, az internethasználók tábora az ezredfordulón félmilliárdosra duzzadt. 2003-ban Amerikában már újra nőni kezdtek a vállalati információtechnológiai kiadások. A technológiai tőzsdeindex a nagy zuhanást követően emelkedni kezdett, de a dotcom válságot követően több hagyományos iparág (pl. energia, nyersanyagok) egy darabig jobb teljesítményt mutatott nála. A feltápászkodó technológiai piacot megtisztította a darwini természetes kiválasztódás, a termékek szabványosodtak, *tömegcikkesedtek*, olcsóbbak lettek, ami kedvező volt a felhasználók számára, és a szolgáltatások felé terelte a versenyzőket¹⁰. A rohamtempóban felépített internetes infrastruktúra sokaknak kiváló lehetőségeket nyújtott a modernizálására, az üzletmenet áramvonalasítására, racionalizálási, kiszervezési akciókra.

A hangsúly a birtoklásról a *használatra* helyeződött át. A lakosság örömmel fogadta az új, hatékony és olcsó kommunikációs lehetőségeket. Az internet valóban elkezdte átalakítani a világot. Témánk szempontjából a legfontosabb kérdés az, hogy milyen hatást gyakorolt a *gazdasági növekedésre* és a *termelékenységre*.

3. A GDP és a termelékenység növekedése a hatvanas évek óta

Először nézzük meg, milyen hatással voltak az *1. sz. táblázatban* felsorolt általános célú technológiák a gazdasági növekedésre és a termelékenységre! A feladat nem egyszerű, az eredményeket óvatosan kell kezelnünk, mivel egy sor mérési problémába ütközünk.

Ha hosszabb időszakra vagyunk kíváncsiak, bele kell nyugodnunk, hogy az adatokban sok a bizonytalanság. Vegyük először a gazdasági növekedés mérésének kérdését. Az 1929-33-as nagy gazdasági világválság idején (ami minden bizonnyal összefüggésben állt az autóipari nagy innovációs hullámmal) az amerikai állam érzékelte, hogy a gazdaság bajban van, tömegek szegényednek el, de azt nem tudta, hogy mekkora volt a gazdasági visszaesés. Ha az államok be akartak avatkozni a gazdaság működésébe, ha valamilyen gazdaságpolitikát szándékoztak érvényesíteni, statisztikai adatokra és mérőszámokra volt szükségük. Az amerikai Simon Kuznets 1934-ben lépett elő a nemzeti jövedelem mérését célzó elképzeléseivel. A GDP (*Gross Domestic Product*, magyarul bruttó hazai termék) mérésének módszertanát csak a második világháború után intézményesítették globálisan.

Az ötvenes évektől a termelékenység (mennyi terméket állítanak elő a munkások egy óra alatt) is nagy figyelmet kapott. Robert Solow 1957-ben publikált modellje lehetővé tette a közgazdászok számára, hogy a tőke, illetve a munkaerő bővülésének köszönhető növekedést módszeresen elválasszák az innováció által okozott (*Total Factor Productivity*, teljes tényezőtermelékenység) gyarapodástól.

Mindezekből az is látható, hogy az *1. sz. táblázatban* szereplő korai általános célú technológiáknak a gazdasági növekedésre gyakorolt hatását csak becsülni lehet; statisztikusok helyett inkább a gazdaságtörténészek következtetéseire kell támaszkodnunk. A GDP-t valóban sokáig nem mérték, de a történészek eléggé megbízható adatokat találtak például a beszedett adókról és vámokról, a

¹⁰ Gerstner, 2002

megművelt földterületek nagyságáról, az éves termés gazdagságáról, és ezekből következtetni tudtak az általános gazdasági teljesítményre is.

Ha a globális GDP-re vonatkozó statisztikai adatokat (az összehasonlíthatóságra vigyázva) kiegészítjük a történeti becslésekkel, érdekes következtetésre juthatunk. A gazdasági teljesítmény a 18. századot megelőzően alacsony volt és alig változott: a GDP növekedését mutató görbe gyakorlatilag lapos volt, a gazdaság stagnált, némi növekedést csak a népesség lassú gyarapodása okozott. Az ipari forradalom beindulásával azonban a helyzet megváltozott, a görbe meredeken, mondhatni exponenciálisan emelkedni kezdett¹¹. Az 1. sz. táblázatban feltüntetett általános célú technológiák, a gyorsan szaporodó innovációk minden bizonnyal a növekedés motorjai voltak, nélkülük ez a drámai változás nem következhetett volna be: a világ belépett a „modern gazdasági növekedés” korszakába.

Az általános kép tehát egyértelmű, a részletek tekintetében azonban már akadnak ellentmondások és nyitott kérdések is.

Nézzük meg először, hogyan növekedett a globális GDP a hatvanas évek eleje óta, ami már több mint fél évszázados időtávot jelent. A Világbanktól átvett adatokat a 2. sz. táblázat tartalmazza. Tekintsünk most el a mérési bizonytalanságoktól, és vizsgáljuk meg, miben különböznek az egyes évtizedek!

2. táblázat. A globális GDP évi átlagos növekedése évtizedek szerinti bontásban

Időszak	GDP növekedése (éves átlag, %)
1960-69	5,4
1970-79	4,1
1980-89	3,0
1990-99	2,7
2000-09	2,9
2010-19	3,2
2020-25	2,6

Forrás: World Bank, 2026a

A hatvanas éveket és a hetvenes évtized elejét a történészek gyakran a kapitalizmus aranykoraként emlegetik. Az éves növekedés átlaga meghaladta az 5%-ot, ami néhány különleges körülménynek volt köszönhető: még tartott a második világháború utáni fellendülés, rengeteg amerikai tőke áramlott Európába és Japánba, a tőke technológiai és vezetési-szervezési innovációkat, modern tudást hozott magával; az olajválság előtt olcsó volt az energia; a foglalkoztatási szerkezetben csökkent az alacsony termelékenységű mezőgazdaság, és nőtt a magas termelékenységű ipar aránya; a háború utáni „baby boom” következményeként sok fiatal és képzett munkás jelent meg a munkaerőpiacon, akik keresni és fogyasztani kezdtek; államközi egyezményeknek köszönhetően lendületesen fejlődött a nemzetközi együttműködés és kereskedelem.

A nyolcvanas és a kilencvenes évtized a globális átalakulás és az egyenlőtlenség időszaka volt. A növekedés egyes országokban lelassult, másokban viszont a technológiai innovációk fellendülést hoztak. A globális átlag, ahogy a 2. sz. táblázaton láthatjuk, elmaradt a hatvanas évek rekordjától. A hetvenes évek vége a *stagfláció* jegyében telt: lassú növekedés, gyorsan emelkedő árak. Az infláció ellen harcolva Amerikában alaposan megemelték a kamatlábakat, ami stabilizálta a dollárt, de globális recessziót okozott, a fejlődő országokra pedig óriási adósságterheket rótt és

¹¹ Our World in Data, 2026

visszafogta a növekedésüket. A politikában Reagan és Thatcher idején a *neoliberalizmus* irányzata érvényesült deregulációval, magánosítással, adócsökkentéssel.

A kilencvenes évek hangulata jobb volt: véget ért a hidegháború, a szovjet blokk megszűnésével emberek százmilliói kapcsolódtak be a piacgazdaságba, megkezdődött Kína és India átalakulása és felemelkedése¹². Megérkezett az internet, a világ összekapcsoltabb, hálózatosabb lett, az üzleti iskolákban mindenfelé az „új gazdaságról”, a szolgáltatások, a szoftver, a számítógépes hálózatok, az olcsó és korlátlan kommunikáció fontosságáról beszéltek. A kor fontos tanulsága, hogy az összekapcsolt világban egy ország válsága azonnal súlyos problémákat okozhat másfelé is.

A számítógépek a hetvenes évektől kezdődően, tehát már évtizedekkel az internetes korszak előtt megjelentek a vállalatoknál is, felmerül tehát a kérdés, hogy a használatuk mennyiben járult hozzá a növekedéshez.

4. A „termelékenységi paradoxon”

Fentebb részletesebben foglalkoztunk az internetes innovációs hullámmal, elemeztük az ezredforduló idején bekövetkezett dotcom válság okait és következményeit. Most tegyük fel újra a kérdést: milyen hatással volt a számítógépek, majd az internet megjelenése és terjedése a gazdasági növekedésre? A válasz nem egyértelmű, a kérdésről ma is viták folynak. Ahogy a fenti rövid és elnagyolt elemzésből is láthattuk, a gazdasági növekedésre számos tényező hat (gazdasági ciklusok, politikai fordulatok, állami gazdaságpolitika stb.), ezért szinte lehetetlen pontosan megmondani, hogy egy-egy konkrét időszakban mekkora a technológiai innováció, az egymást váltó, egymásra torlódó általános célú technológiák hatása.

Az informatikai befektetések és a termelékenység közötti ellentmondásos kapcsolatot a szakirodalom „termelékenységi paradoxonnak” nevezi. A Nobel-díjas Robert Solow-nak tulajdonítják azt a kijelentést, miszerint „a számítógépek mindenütt megtalálhatók, kivéve a gazdasági statisztikákat”¹³. A tudós az USA gazdasági adataira hivatkozott: a termelékenység növekedése a hetvenes és a nyolcvanas években a korábbinál alacsonyabb volt (lásd a 3. sz. táblázatot), miközben az informatikai beruházások lendületesen növekedtek.

3. táblázat. A termelékenység évi átlagos növekedése az Amerikai Egyesült Államokban (*nonfarm business sector*)

¹² Lásd pl. Das, 2001; Kroeber, 2016

¹³ Solow, 1987

Időszak	Évi átlagos növekedés (%)
1947 Q1 – 1973 Q4	2,7
1973 Q4 – 1980 Q1	1,4
1980 Q1 – 1990 Q3	1,6
1990 Q3 – 2001 Q1	2,1
2001 Q1 – 2007 Q4	2,8
2007 Q4 – 2019 Q4	1,5
2019 Q4 – 2024 Q3	1,8
<i>Teljes időszak (1947-2024)</i>	<i>2,1</i>

Forrás: U.S. Bureau of Labor Statistics, 2026

Lester Thurow, az MIT üzleti iskolájának dékánja megállapította, hogy 1978 és 1985 között az USA teljes gazdasági outputja 15%-kal bővült, a fizikai munkások száma pedig 6%-kal csökkent; mindeközben azonban az irodai dolgozók (a menedzsereket is ideértve) tömege 21%-kal gyarapodott, ellensúlyozva a munkások termelékenységének javulását, és összességében rontva az ország gazdasági versenyképességét. A „termelékenységi pangás” időszaka egybeesett a nagy irodaautomatizálási akciókkal, amelyeknek éppenséggel javítaniuk kellett volna a fehérgallérosok teljesítményét¹⁴.

Ezek a vélemények, és az általuk gerjesztett viták arra is rávilágítanak, hogy az informatika gazdasági szerepét, üzleti értékét korántsem látjuk tisztán se makrogazdasági, se vállalati szinten.

Egyesek (például a nemrég elhunyt Alan Greenspan¹⁵) a paradoxon okát a mérések pontatlanságában, módszertani gyengeségekben keresték. A jól mérhető ipari korszak véget ért, mondták, a szolgáltatások növekedését pedig nehezen lehet mérni hagyományos statisztikai eszközökkel. A mérhető mennyiségnél különben is fontosabb a sokkal nehezebben megragadható minőség, a termékek és szolgáltatások változatossága, a testre szabás, a kényelem, az általános jólét. Mások szerint az említett időszakokban a termelékenység különböző okok miatt éppenséggel csökkent volna, és pont az informatikának köszönhető, hogy ez nem következett be; az informatika hatása tehát pozitív, csak más erők ellensúlyozzák. Ismét mások a technika gyermekbetegségeire hivatkoztak: a sokféle gép és rendszer, az új technológia természetes gyengeségei sok problémát okoznak, ami leköti az emberek energiáját – ez a helyzet azonban átmeneti, idővel az alkalmazás és az integráció egyre könnyebb lesz, a termékek pedig egyre jobbak.

A 3. sz. táblázat adataiból mindenesetre arra következtethetünk, hogy a technológiai innovációban élen járó Amerikai Egyesült Államokban az 1973 előtti „aranykor” után hosszabb ideig alacsony volt a termelékenység növekedése, pedig pont ekkor jelentek meg a vállalatoknál a mainframe számítógépek. Fordulat csak a kilencvenes években és az új évszázad elején következett be, és csak néhány évig tartott.

¹⁴ Thurow, 1986

¹⁵ Greenspan, 2007

5. Fellendülés a dotcom válság előtt és után

A 3. sz. táblázatban található adatokból jól látszik, hogy az Amerikai Egyesült Államokban a kilencvenes években a megelőző időszakhoz képest gyorsabban nőtt a termelékenység. A 4. sz. táblázat a kilencvenes éveket két részre bontja, amelyek között egyes kiemelt makrogazdasági mutatók tekintetében markáns különbségek mutatkoznak. A múlt század vége az internet berobbanásának és általános elterjedésének időszaka volt. Fontos határkőnek számít a bárki által könnyen használható Navigator böngészőt fejlesztő Netscape tőzsdei megjelenése 1995 augusztusában. A kibontakozó internetes lázat, a pénzügyi lufi felfúvódását és kidurranását fentebb már leírtuk; a következőkben az új általános célú technológia gazdasági hatásaira koncentrálnunk.

4. táblázat. Gazdasági mutatók a Netscape tőzsdei megjelenése előtt és után, USA

Gazdasági mutató	1991 Q1 – 1995 Q3	1995 Q3 – 2000 Q1
A GDP évi átlagos növekedése (%)	3,0	4,3
A termelékenység évi átlagos növekedése (%)	1,7	2,8
A munkanélküliség évi átlagos rátája (%)	6,6	4,8
Évi átlagos maginflációs ráta (%)	3,3	2,3

Forrás: Mandel, 2000, 15. o. Az U.S. Bureau of Labor Statistics és Bureau of Economic Analysis adatai alapján.

Nézzük meg először, hogy mi történt a kilencvenes évek elején. Az USA gazdaságában 1990 júliusában recesszió kezdődött, ami 1991 márciusáig tartott. A visszaesés enyhe volt és rövid, nyolc hónap után újra megindult a növekedés. A recessziót követő két-három évet azonban „jobless recovery”-ként emlegetik a gazdaságtörténészek. A recessziók elmúltával a vállalatok általában toborzásba kezdenek, a munkanélküliségi mutató csökkenni kezd. Most azonban más volt a helyzet: a recesszió 1991-ben „hivatalosan” véget ért, de a magas munkanélküliségi ráta megmaradt: 1992 júniusában tetőzött 7,8%-kal.

E szokatlan jelenségnek több oka volt. Tudjuk például, hogy a hidegháború végét elhozó globális politikai fordulatok következtében az USA visszafogta a katonai költségvetését, a csökkenő kereslet pedig elbocsátásokkal is járt. Ugyanakkor okunk van feltételezni, hogy a foglalkoztatási gondok a technológiai fejlődéssel is összefüggtek. A vállalatok, különösen a legnagyobbak, felismerték, hogy infokommunikációs eszközeik és rendszereik segítségével számos feladatot olcsóbban, gyorsabban, kevesebb alkalmazottal is meg tudnak oldani.

Látványos átszervezési, automatizálási programok indultak sokfelé, ez volt a *Business Process Reengineering*, a vállalati folyamatok radiális újraszervezésének virágkora. Nem meglepő, hogy a leépítések elsősorban az irodai dolgozókat és a középszintűket érintették, őket lehetett leginkább számítógépekkel kiváltani. A vállalatok, köztük ismert óriások is, tízezzel bocsátották el az embereket, ami szokás szerint politikai következményekkel is járt: a munkanélküliség kényes kérdésként jelent meg az 1992-es elnökválasztás idején. Felélénkült a vita a technológiai fejlődés munkaerőpiaci következményeiről: valóban elvehetik a számítógépek az emberek munkáját?

A 4. sz. táblázaton láthatjuk, hogy a kilencvenes évtized második felében már alacsonyabb volt az átlagos munkanélküliségi ráta, mint az elsőben. Az internetes láz idején (a lázat a nagy elbocsátásokkal megtakarított, befektetési lehetőségeket kereső pénz is fűtötte) a beruházások, a gyorsan növekvő vállalkozások felszívták a munkaerőt. A GDP gyorsan nőtt, a termelékenység látványosan emelkedett, és az infláció sem okozott gondot. Könnyen adódik a következtetés: az internet általános célú technológiaként jó hatással volt az amerikai gazdaságra, a növekedésre és a termelékenységre egyaránt.

Az általános kép tehát pozitív, a részletek tekintetében azonban nem árt az óvatosság. A McKinsey Global Institute részletes elemzést tett közzé az információs technológia és a termelékenység kapcsolatáról¹⁶. Megállapították, hogy valóban nőtt a termelékenység, de ez a növekedés néhány iparágra koncentrált, leginkább magára az infokommunikációs szektorra. A termelékenység javulásának 76%-a hat gazdasági szektor (az említett infokommunikáció alszektorai mellett a kis- és nagykereskedelem, valamint az értékpapírok forgalmazása) teljesítményével volt magyarázható, amelyek együttesen alig egyharmadát adták a GDP-nek. A technológiai hatások mellett azt is számításba kell venni, hogy a láz idején kibontakozó heves verseny hatékonyságot növelő lépésekre ösztönözte a vezetőket, és ez nem jelentett mindenütt informatikai fejlesztéseket.

Mindezek fényében pontosítsuk a fenti következtetést: az internet az adott korban jó hatással volt a gazdasági teljesítményre, de nem egyszerre és nem mindenütt. Egyes iparágakban valóban nagy volt a szerepe, másokban viszont elenyésző. A tisztánlátást az is nehezíti, hogy a termelékenység szempontjából sikeres szektorokhoz (így például a kereskedelemhez) tartozó vezető cégeknél az informatikai fejlesztések mellett egyéb átalakítási programok is folytak, vagyis az IT hatását nehéz kibontani a tarka képből. Ráadásul arra sincs garancia, hogy az informatika által biztosított értéknövekedés a beruházóknál jelenik meg nagyobb termelékenység vagy növekvő profit formájában: több példát is felhozhatunk arra, hogy valamilyen technológiai fejlesztéssel a vevők járnak jól, a fejlesztők profitját viszont elolvasztja a verseny. A fejlesztések akár túlkínálatot is előidézhetnek, mire az árak csökkenni kezdenek, a pénzügyi mutatók romlanak. Az informatikai befektetések önmagukban nem hoznak jelentős üzleti eredményeket, a technikai fejlesztési akciók csak tágabb stratégia-szervezeti összefüggésrendszerben értékelhetők.

Térjünk most vissza a 3. sz. táblázathoz. Fentebb bemutattuk, hogy az internetes láz az ezredforduló idején a technológiai lufi kipukkanásával végződött. A vállalatok visszafogták az informatikai kiadásait, a kereslet csökkent. A táblázaton ugyanakkor azt látjuk, hogy a krízis után a termelékenység növekedése nem lassult le a gazdaságban, a lendület megmaradt. Úgy tűnik számolni kell *késleltetett hatással* is. Erik Brynjolfsson és Adam Saunders (mindketten neves egyetemek professzorai) pontosan erre a következtetésre jutottak: vállalati adatokat elemezve megállapították, hogy rövid távon a számítógépek hozzájárulnak az output növekedéséhez, de a hatásuk csak idővel bontakozik ki, a termelékenység javulására csak évek múlva lehet számítani. Ami még fontosabb: az informatikai eszközök beszerzése és hadrendbe állítása nem elég, intézményi változásokra, átszervezésekre is szükség van, de ezek nem mennek egyik napról a másikra, amit más iparágak és innovációk története is igazol¹⁷. Ha ez az állítás igaz, az USA gazdaságában az ezredfordulón bekövetkezett termelékenységi ugrás a korábbi nagy összegű informatikai beruházásoknak volt köszönhető, vagy azoknak is.

Úgy is fogalmazhatunk, hogy információs technológia *szükséges, de nem elégséges* feltétele lehet a termelékenység növekedésének. Az informatika lehetőségeket teremt, amiket ki is kell használni, például át kell szervezni a vállalati folyamatokat, át kell képezni a munkásokat, új termékekkel, szolgáltatásokkal és üzleti modellekkel kell megjeleníteni.

Erik Brynjolfsson és munkatársai a késleltetett hatások jelenségét „termelékenységi J-görbének” nevezik, ahol a J a termelékenység változását mutató görbe alakjára utal¹⁸. Az elgondolás lényege az, hogy a jelentős általános célú technológiák hasznosításához komplementer, vagyis kiegészítő jellegű beruházások is szükségesek. Nem elég számítógépet és szoftvereket vásárolni: a folyamatokat át kell szervezni, az alkalmazottakat át kell képezni, mindenféle fejlesztési programokat kell indítani. Ezek nagyrészt *immateriális beruházások*, amelyek a vállalat *tudástőkéjét* növelik, de nem produkálnak azonnali hasznot, amit a statisztikusok mérni tudnának.

¹⁶ Idézi Carr, 2004

¹⁷ Brynjolfsson, Saunders, 2010

¹⁸ Brynjolfsson, Rock, Syverson, 2021

Mivel a vállalat úgy költ rájuk, hogy közben nem nő a bevétele, úgy tűnik, hogy a termelékenysége csökken. (Ez a völgy a J elején). Lehet, hogy az immateriális javak, a tudástőke csak évek múlva fordulnak termőre: a befektetési szakasz lezárul, eljön az aratás ideje, a vállalat alkalmazottai több értéket állítanak elő, a termelékenység mutatója nőni kezd a hivatalos statisztikákban.

Nem kizárt, hogy a mesterséges intelligencia gazdasági hatása is a J-görbe mentén fog kibontakozni. Az sem kizárt, hogy a közeli jövőben a dotcom válságra emlékeztető krízist is át kell élnünk.

6. Korunk általános célú technológiája: a mesterséges intelligencia

Az elmondottak tükrében vizsgáljuk meg, milyen hatást gyakorolhat korunk legfontosabb általános célú technológiája, a mesterséges intelligencia a gazdaságra, ismét a gazdasági növekedést és a termelékenységet állítva a középpontba. Leginkább arra vagyunk kíváncsiak, hogy megismétlődnek-e most is az internetes hullám idején, vagyis a kilencvenes és a kétezres évek elején tapasztalt jelenségek, vannak-e J-görbére utaló jelek. Természetesen figyelembe kell venni azt is, hogy az MI hullám egyáltalán nem tekinthető lefutottnak, 2026 elején, amikor a jelen írás született, még egyértelműen az emelkedő fázisban járt.

A mesterséges intelligencia esetében a lappangási időszak hosszú volt, több mint fél évszázadig tartott, hiszen az elnevezés az ötvenes években jelent meg. Az első évtizedek MI fejlesztőit az a meggyőződés vezérelte, hogy az emberi tudás absztrakt jelekkel és matematikai műveletekkel kódolható, és ilyen formában számítógépbe táplálható. Az elképzelés szerint a gép hamarosan alkalmas lesz arra, hogy kérdésekre válaszoljon, szövegeket fordítson, képeket ismerjen fel. Az így előállított „szakértői rendszerek” különösen népszerűek voltak a nyolcvanas években, fejlesztésükhöz elsősorban a formális logikában gyökerező Prolog programozási nyelvet használták.

Bár születtek figyelemre méltó eredmények, készültek hasznos rendszerek, kiderült, hogy ez az irány zsákutca, a gép nem birkózik meg az emberi intelligenciát igénylő feladatok komplexitásával, és a szakértők sem lelkesedtek azért, hogy a tudásukat gépre vigyék.

Az új évszázad elején a *gépi tanulás* került a figyelem középpontjába. A filozófia megváltozott: nem az emberi tudást kell programba írni, hanem rétegekbe szervezett mesterséges neuronokból álló neurális hálókat és tanuló-optimalizáló algoritmusokat kell használni. A gépet példákkal kell tréningezni, „betanítani”, ami így tulajdonképpen önmaga fejleszti a döntési modelljét. A *mélytanulás* a tízes évek elején mutatta meg a nagyvilágnak oroszlánkörmeit, amikor egy torontói számítógépes képfelismerési versenyen egy gépi tanulással tréningezett modell kiváló teljesítményt nyújtott, a gép később ezen a területen az emberek teljesítményét is túlszárnyalta¹⁹.

A vállalatok érdeklődését gyorsan felkeltette az új megoldás. A nagy tőzsdei cégek 2010 környékén még meg sem említették a negyedéves beszámolóikban a mesterséges intelligenciát, 2017-ben viszont már több százan jelezték, hogy érdekli őket és valamit akarnak vele kezdeni. A lázat tovább fokozta a nagy nyelvi modellek megjelenése: a ChatGPT 2022 novemberében vált a nagyközönség körében elérhetővé, felhasználói tábor a pillanatok alatt sok millióra duzzadt. Gyorsan szaporodtak az MI vállalkozások, a kockázati tőke gyorsan rárepült a lehetőségekre.

Nézzük meg, hogy mi történt mindeközben a gazdaságban!

A 2. sz. táblázathoz visszatérve láthatjuk, hogy évszázadunk tízes éveiben a globális GDP átlagos éves növekedése valamivel magasabb volt, mint az ezredforduló utáni évtizedben. Az utóbbi

¹⁹ Bögel, Nagy, 2025, 2. fej.

időszak a dotcom válsággal kezdődött, ami után a gazdaság hamar magához tért és szép GDP növekedési mutatókat produkált, a lendületet viszont megtörte az évtized végének súlyos pénzügyi világválsága. Az amerikai gazdaság recessziója 19 hónapon át tartott, hivatalosan 2009 júniusában fejeződött be. A válságot a szokásos visszapattanás követte, majd egy viszonylag nyugodt időszak következett mérsékelt növekedéssel és alacsony inflációval. Nem következett be olyan gazdasági robbanás, mint a dotcom válság után, de az újabb kríziseket egészen a tízes évek végéig sikerült elkerülni. Bár több nyugati országban stagnáltak a hagyományos gyártási szektorok, a digitális gazdaság építése lendületesen folyt, látványosan fejlődtek azok a technológiák és termékek (pl. felhő-számítástechnika, Big Data, okostelefonok), amelyek a mai mesterséges intelligenciához kellenek²⁰.

2019-ben kitört a koronavírus-világjárvány, ami 2020-ban a bezárások miatt recessziót hozott a világgazdaságban²¹. 2021-ben a világ visszazökkent a normális kerékvágásba, a globális GDP növekedés meghaladta a 6%-ot, ami az alacsony bázis mellett minden bizonnyal a határozott kormányzati intézkedéseknek, a gazdaságba pumpált támogatásoknak is köszönhető. Ezt követően az éves növekedési mutató egészen napjainkig beállt a 2,6-2,7%-os szintre, ami az elmúlt fél évszázadot tekintve gyenge eredménynek számít.

A Világbank 2026 januárjában kiadott jelentése²² érdekes ellentmondásokra hívja fel a figyelmet. Megállapítja, hogy a sorozatos sokkok (pénzügyi válság, világjárvány, magas infláció a covidot követően, háború, vámháború, váratlan politikai irányváltások) ellenére a gazdaság meglepően reziliensnek bizonyult: talpra tudott állni a csapások után. A jelentés a 2025 évi globális GDP növekedést 2,7%-ra becsüli. Mindezek ellenére aggodalomra is van ok. A jelentés szerint, ha a Világbank több éves növekedési előrejelzése igaznak bizonyul, a húszas évtizedben lesz a legalacsonyabb a növekedés a hatvanas évek óta. „Ez nem az a dagály, amely minden csónakot megemel” – állítja a dokumentum²³. A fejlett országok gazdagabbak lesznek, a fejlődők közül viszont sokan nehéz helyzetbe kerülnek, különösen azok, amelyek súlyos gazdaságpolitikai hibákat követtek el, és most nem tudnak korrigálni. A világjárvány utáni talpra állás igazából nem a gazdaság erejének köszönhető, hanem radikális, nehezen megismételhető manővereknek. A prognosztizált növekedés egy sor országban kevés lesz a szegénység visszaszorításához, elfogadható munkalehetőségek biztosításához: mélyülni fog a szakadék a gazdag és az elmaradott országok között.

A 3. sz. táblázaton láthattuk, hogy az új évszázad első éveiben az Amerikai Egyesült Államokban a termelékenységi éves átlagos növekedése rekordmagasságot, 2,8%-ot ért el. Az U.S. Bureau of Labor Statistics (BLS) fentebb már hivatkozott számításai szerint ugyanez a mutató 2019 vége és 2025 harmadik negyedéve között 2,0% volt (nagy kilengések mellett), ami nagyjából annyi, mint az 1947 utáni átlag. Figyelemre méltó ugyanakkor, hogy a hivatal 2026 január 29-én kiadott jelentése szerint a „nonfarm”, vagyis a mezőgazdaságon kívüli üzleti szektorban 2025 harmadik negyedévében 4,9% volt a munka termelékenységi növekedése az előző negyedévhez képest, vagyis a növekedési görbe 2025-ben elindult felfelé, amiből még nyilván nem szabad merész következtetéseket levonni az MI hatására vonatkozóan.

A BLS az innováció hatását pontosabban mutató teljes tényező-termelékenységi (TFP) változásáról is közöl adatokat. Számításai szerint a mezőgazdaságon kívüli üzleti szektorban a 2019 és 2024 közötti időszakban éves átlagban 0,9% volt a TFP növekedése, ami nagyjából annyi, mint az 1990 és 2000 között mért mutató, és több, mint a 2007 és 2019 közötti 0,6%. A TFP 2000 és 2007 között emelkedett a leggyorsabban, amikor a gazdaságban az output növekedése jelentősen meghaladta a kombinált inputokét, jelezve, hogy a vállalatok növelni tudták a

²⁰ Bögel, 2015

²¹ Mátyás, 2022

²² World Bank, 2026b

²³ i.m. xvi old.

hatékonyságukat. Lehet, hogy most is valami hasonló történik, és ebben már a mesterséges intelligenciának is szerepe van.

Az átlagok eltakarják az eredmények szóródását. Javuló teljesítmény elsősorban a csúcstechnológia és a tőkeigényes szektorokban figyelhető meg. A vezető szektorok közé tartoznak az IT, a pénzügyek és a technikai és szakértői szolgáltatások; a kiskereskedelem és a tartós fogyasztási cikkek gyártása vegyes képet mutat; a lemaradók körében jelenik meg a bányászat, az egészségügy és a személyi szolgáltatások csoportja. Az utóbbiakban hagyományosan nehéz növelni az egy főre eső outputot.

Figyelemre méltó jelenség, hogy miközben Amerikában emelkedik a termelékenység, a munkaerő részesedése a nemzeti jövedelemből 2025-ben rekordmélységben járt. Ha a termelékenység javulását az MI használatának és az automatizálásnak tulajdonítjuk, azt is el kell ismernünk, hogy a keletkezett haszon a tőke tulajdonosainál és az óriási IT cégeknél csapódott le.

Mi a helyzet Európában? A második világháború után a vén kontinensen nagyobb volt a munka termelékenységének éves növekedése, mint Amerikában, ami leginkább a Marshall-segélynek, a pusztulás utáni sikeres talpra állásnak, integrációs lépéseknek, az okos gazdaságpolitikának volt köszönhető. A lendület azonban a hetvenes évekre kimerült, egyre több gazdaságszerkezeti probléma jelentkezett. Európa csak lassan fogadta be az infokommunikációs innovációkat. A termelékenység növekedése lelassult, az USA a tőkeerejével, saját óriási és egységes piacával elhúzott az innovációs versenyben. A képen csak csak egyes jobban teljesítő országok, különleges innováció „szigetek” javítanak valamelyest. Európa versenyképességi problémáival és az MI jelentőségével részletesen foglalkozik a Draghi-jelentés²⁴. Ha a helyzet nem változik, félő, hogy az USA és az EU közötti szakadék mélyülni fog. Friss adatokból látható²⁵, hogy az előbbiben 2020 és 2024 között ötször annyi magántőke áramlott az MI szektorba, mint az utóbbiban. (A magánbefektetések nagysága tekintetében Európa kicsivel megelőzte Kínát, ott viszont az állami pénzeket is számításba kell venni.)

Visszatérve a globális gazdasági színtérre: egyes biztató, de még halvány jelek ellenére a gazdaságban nem láthatunk figyelemre méltó fellendülést, radikális fordulatot. A GDP növekedésének mutatói általában alacsonyok, és a munkanélküliség is normálisnak tekinthető (illusztrációképpen néhány ország adatait megjelenítjük az 5. sz. táblázatban.)

5. táblázat. A GDP 2025. évi növekedése és a munkanélküliségi ráta egyes országokban

²⁴ European Commission, 2024

²⁵ Stanford University, 2026

Ország	GDP növekedése (%)	Munkanélküliségi ráta (%)
USA	2,1	4,4
Kína	5,0	5,1
India	7,4	6,9
Oroszország	0,6	2,1
Egyesült Királyság	1,4	5,1
Euró övezet	1,4	6,3
Ausztria	0,6	5,8
Németország	0,2	3,8
Hollandia	1,5	4,0
Csehország	2,5	2,8
Dánia	2,8	2,9
Malajzia	4,9	2,9
Dél-Korea	1,3	4,1

Forrás: The Economist, 2026. jan. 24, 73. o. A GDP adat a The Economist Intelligence Unit becslése.

Az egy főre eső MI beruházások nagyságát tekintve jelentős különbségeket láthatunk egyes országok között. A magánbefektetések tekintetében a mezőnyt az USA vezeti, de köztudott, hogy vannak lélekszámukat tekintve kicsi, de igen „MI intenzív” országok, mint például Szingapúr, ahol az egy főre jutó befektetés majdnem annyi, mint Amerikában. Kínában az állam óriási összegeket fordít MI fejlesztésekre és beruházásokra. Az MI fejlődése szempontjából nagyon fontosak a nagy adatközpontok: ezek közel 40%-a jelenleg Amerikában van. Egyes országokban rengeteg pénz áll rendelkezésre MI kutatásra és fejlesztésre, de sok olyan hely van, ahol ez az összeg elenyésző méretű.

Nagy különbségeket, mélyülő szakadékokat jelez a Microsoft 2025. évi MI diffúziós jelentése²⁶ is. Az elemzés megállapítja, hogy az új eszközök (leginkább a generatív AI) használatának terjedése a Globális Észak országaiiban közel kétszer olyan gyors, mint a Globális Délen: az előbbieken munkaképes korú lakosság 24,7%-a használja már ezeket az eszközöket, míg az utóbbiaknál csak 14,1%. Az adott év végén a diffúziós rangsort az Emírátságok, Szingapúr, Norvégia, Írország és Franciaország vezette, Magyarország a 19. helyen állt megelőzve például az Egyesült Államokat, Csehországot és Finnországot is. (Ajánlatos ezeket a helyezéseket más felmérések eredményeivel összevetni.) Úgy tűnik, hogy az IT fejlesztésekbe már régóta nagy összegeket fektető országok az MI korában is meg tudták őrizni a vezető helyüket, az összehasonlításoknál azonban ajánlatos óvatosan eljárni.

7. Következtetések és kérdőjelek

Milyen kép bontakozik ki mindezekből? Milyen következtetéseket vonhatunk le a gazdasági növekedés és a mesterséges intelligencia innovációs hullámának kapcsolatáról?

Fenntartások nélkül kijelenthetjük, hogy az infokommunikációs szektor a modern gazdaság masszív eleme. A tágabb értelemben vett digitális gazdaság (ami például az online kereskedelmet is magába foglalja) hozzájárulása a globális GDP-hez becslések szerint 15-17%-ra tehető²⁷. Növekedése jelentősen meghaladja a globális gazdaság egészét, vagyis a globális GDP gyarapodása jelentős részben a digitális szektor tempójának tulajdonítható. Más szektorok növekedésére is jó

²⁶ Microsoft, 2025

²⁷ World Bank, 2025

hatással van: a digitalizáltság általában magasabb termelékenységgel jár együtt. Friss adatokból az is látható, hogy napjainkban a növekedés első számú motorja a mesterséges intelligencia, ami már több mint a felét teszi ki a technológiai kiadásoknak.

Tételezzük fel, hogy a mesterséges intelligencia innováció ciklusa olyan lesz, mint amit fentebb modellként bemutatunk, és az internetes hullám példájával illusztráltunk. Hol járhat most az MI az maga innovációs hullámában? A lappangáson és a berobbanáson minden bizonnyal túl. A jelen kérdése az, hogy számíthatunk-e egy olyan innovációs lufira, mint amilyet az internet esetében a kilencvenes években és az ezredfordulón megfigyelhattunk. Nem biztos, hogy ez a jelenség most is bekövetkezik, de van rá esély. Lássuk, hogy milyen jelekre érdemes odafigyelni!

Az MI beruházások kétségtelenül óriásira nőttek az elmúlt években. Az UBS bankház számításai szerint²⁸ ha a hat közismert technológiai óriáscég mellé odateszünk még három élvonalbeli MI stratupot, láthatjuk, hogy az összesített beruházásaik 2017 és 2025 között több mint a hétszeresükre emelkedtek.

2025 első felében az amerikai GDP növekedését az IT beruházásokban bekövetkezett boom adta. Így is mondhatjuk: e nélkül nem lett volna növekedés. A tőzsde értékének emelkedése fele részben tíz vállalatnak (Nvidia, Amazon, Meta stb.) volt köszönhető. A nyugati világ kockázati tőkéjének harmadát MI cégekbe fektették. Van miből választaniuk: rengeteg MI vállalkozás jelent meg világszerte.

Az MI cégek értékelésénél ugyanazt a jelenséget figyelhetjük meg, mint az ezredforduló előtt az internetes vállalkozásoknál. Az MI világ két vezető modellfejlesztője, az OpenAI és az Anthropic milliárdos befektetéseket hajtanak fel hónapról hónapra, együttes piaci értékük 2025 végén megközelítette a féltrillió dollárt. Őket már tulajdonképpen unikornisoknak sem lehet nevezni, új elnevezéseket kell keresni a zoológiai enciklopédiákban.

Akárcsak a korábbi lufiknál, most is van szerepe a csodavárásnak, hiszen az MI „szent grálja”, az *általános mesterséges intelligencia* (angol rövidítéssel: AGI) még nem született meg, de többen bábáskodnak a még üres bölcsőnél²⁹. Hogy pontosabbak legyünk, több ilyen bölcső is van, és akinél először felsír az újszülött, az bizony nagyon sok pénzt kereshet. Miért ne lenne érdemes ekkora nyereségért kockázatot vállalni?

A bevételek és a nyereségek egyelőre nem igazolják a különleges érdeklődést és a hatalmas befektetéseket. Az UBS számításai szerint a Nyugat vezető MI cégei jelenleg együttesen évi 50 milliárd dollár bevételt produkálnak. Ez nagy összegnek tűnik, a növekedése is gyors, de kevés a megkezdett, illetve tervezett, több trillió dollárra rúgó adatközpont-beruházások igazolásához. Egyelőre az sem látszik világosan, hogy a bevételekből miként lesz nyereség. Egy MIT-s tanulmány³⁰ szerint a felhasználó szervezetek 95%-a egyelőre semmi hasznót nem lát a generatív MI-ből.

A jelenlegi piaci adatok és a személyes tapasztalatok azt mutatják, hogy az MI eszközökkel, különösen az ingyenesekkel rengetegen játszanak, kísérleteznek: az internet tele van mindenféle vicces képekkel és videókkal, világsztárokkal készített ál-szelfikkel, mókás újévi üdvözlőlapokkal; tanárok számolnak be arról, hogy a tanítványaik közül sokan MI eszközökkel oldják meg a feladataikat; ügyfélszolgálati megkeresésekre sokféle robotok válaszolnak váltakozó sikerrel. Az is kétségtelen, hogy legálisan vagy illegálisan sokan használják az MI eszközöket saját munkájukhoz is. A komoly vállalati alkalmazások száma viszont meglehetősen csekély (a felmérési eredmények gyakran homályosak és ellentmondásosak, például azt sem egyértelmű,

²⁸ Bee, 2025

²⁹ Lásd pl. Ford, 2018; Olson, 2025

³⁰ Wheeler, 2025

hogy ki mit tekint mesterséges intelligenciának), így nem lehet tudni, hogy a kísérletezés, az önálló, illetve a gyakran spontán, központi irányítás nélkül szerveződő csoportos tanulás mikor hoz tényleges eredményeket és gerjeszt masszív vállalati keresletet, amire bizony nagy szüksége lenne az MI-t fejlesztő cégeknek.

A lufit a politika is fújja: egyes államok egymásra licitálva hirdetnek meg nemzeti MI stratégiákat, befektetési, támogatási és deregulációs programokat, fokozva a vállalkozások (leginkább a győztesnek központilag „kijelöltek”) biztonságérzetét: ha baj lenne, az állam majd kisegít bennünket. A tapasztalatok azt mutatják, hogy a politikusok által fújta lufik bizony nagy károkat tudnak okozni.

Kockázatot jelent az is, hogy az MI láz alatt kiépülő kapacitások valószínűleg csak rövid ideig lesznek modernnek. A 19. századi vasútépítési láz alatt épített rengeteg pálya, állomás és mindenféle berendezés sokáig használatban maradhatott; a most épülő adatközpontok, a napjainkban fejlesztett szuperchipek viszont nagy valószínűséggel gyorsan elavulnak, hiszen a fejlődés elképesztően gyors. Az új infrastruktúra hasznos élettartama aligha lesz olyan hosszú, mint a vasúti és az elektromos hálózatoké, vagy akár az ezredfordulón kiépült távközlési kapacitásoké.

Az sem mindegy, hogy kidurranás esetén kiknek kell elviselni a veszteséget. A dotcom lufi esetében a károk sok kisbefektető között oszlottak el. A 19. század közepén begerjedt angliai vasútépítési lázat viszont főleg bankok finanszírozták: a sok bedőlt hitel miatt vissza kellett fogniuk a hitelezést, ami nagy kárt okozott a gazdaságban és elmélyítette a recessziót. A mostani adatközpont-boomot fele részben technológiai óriáscégek pénzelik, méghozzá a saját vagyonukból. Őket aligha sújtaná a földre az MI lufi kipukkanása, bár például az OpenAI-nak ebből a szempontból van félnivalója. A vastag pénztárcájú befektetési alapok miatt sem kell különösebben aggódni. A részvények manapság is jelentős részt képviselnek az amerikai háztartások vagyonaiban, de leginkább a gazdagokéban.

Mi következik mindezekből? Egy MI lufi felfújódásának jelenleg kétségtelenül számos jele van. Hasonló jelenségeket tapasztalhatunk, mint az itt részletesebben bemutatott internetes láz, illetve a korábbi mániák idején: látványos innovációk, túlhevült befektetői és vállalkozói világ, irreális értékelések, homályos megtérülés, kiforratlan üzleti modellek, forró téglák dobálása, csodavárás, nagy ígéretek, szenzációkeltő megnyilatkozások, és így tovább.

Ha a lufi kidurran, sokan bajba kerülhetnek, de egyrészt nem biztos, hogy kidurran, másrészt ha ez mégis bekövetkezne, talán nem dönti az egész gazdaságot recesszióba, a darwini rendrakás még jót is tehet a világnak. Ahogy századunk elején megfigyelhettük, a láz elmúltával egy nyugodtabb, gyümölcsözőbb időszak jöhet, a J-görbe felszálló ágába kerülhetünk, amikor a mesterséges intelligencia valóban megmutathatja az erejét. A mostani makrogazdasági mutatókon a hatása egyelőre nem látszik, de egyes szektorokban biztató fejleményeket figyelhetünk meg.

A korábbi nagy innovációs hullámok tapasztalataiból azt szűrhetjük le, hogy a pozitív hatások kibontakozása nem megy egyik napról a másikra. Az új technológia önmagában nem elég. *egy sor más feltételt is meg kell teremteni.* A GDP és a termelékenység növekedésére a globális gazdaságnak és az Európai Uniónak egyaránt nagy szüksége lenne, akárcsak a több éve kiábrándító eredményeket produkáló Magyarországnak.

Irodalomjegyzék

- Bee, A. (2025): *Are we facing an AI bubble?* UBS, szept. 4. <https://www.ubs.com/global/en/wealthmanagement/insights/marketnews/article.2581859.html>. (Letöltve: 2026. feb. 2.).
- Bögel Gy. (2002): *Az infokommunikációs hullám sajátosságai*. Competitio, szeptember, 73-95.
- Bögel Gy. (2015): *A Big Data ökoszisztémája*. Typotex Kiadó, Budapest.
- Bögel Gy.; Nagy D. A. (2025): *Mesterséges intelligencia a gazdaságban*. Typotex Kiadó, Budapest.
- Brynjolfsson, E.; Saunders, A. (2010): *Wired for Innovation: How Information Technology Is Reshaping the Economy*. The MIT Press, Cambridge, MA
- Brynjolfsson, E.; Rock, D.; Syverson, C. (2021): *The Productivity J-Curve: How Intangibles Complement General Purpose Technologies*. American Economic Journal: Macroeconomics, vol. 13, no. 1, January, 333-72.
- Carr, N. (2004): *Does IT Matter?* Harvard Business School Press, Boston.
- Das, G. (2001): *India Unbound*. Alfred A. Knopf, New York.
- European Commission (2024): *The Draghi report on EU competitiveness*. https://commission.europa.eu/topics/competitiveness/draghi-report_en#paragraph_47059. (Letöltve: 2026. jan. 28.).
- Ford, M. (2018): *Architects of Intelligence: The truth about AI from the people building it*. Packt Publishing, Birmingham.
- Gerstner, L. (2002): *Who Says Elephants Can't Dance?* HarperBusiness, New York.
- Greenspan, A. (2007): *The Age of Turbulence*. Allen Lane, Penguin Books, New York.
- Karikó K. (2023): *Áttörések. Életem és a tudomány*. Helikon, Budapest.
- Kindleberger, C. (2000): *Maniacs, Panics, and Crashes*. John Wiley & Sons, Hoboken, New Jersey.
- Komáromi Gy. (2006): *Anatomy of Stock Market Bubbles*. The ICFAI University Press. Hyderabad.
- Kroeber, A. R. (2016): *China's Economy: What Everyone Needs to Know*. Oxford University Press, Oxford.
- Mandel, M. (2000): *The Coming Internet Depression*. Basic Books, New York.
- Mátyás L. szerk. (2022): *Emerging European Economies after the Pandemic*. Springer, Berlin.
- Microsoft (2025): *Global AI Adoption in 2025. A Widening Digital Divide*. <https://www.microsoft.com/en-us/corporate-responsibility/topics/ai-economy-institute/reports/global-ai-adoption-2025/>. (Letöltve: 2026. feb. 1.).
- Moore, G. (1965): *Cramming More Components onto Integrated Circuits*. Electronics, April 19.
- Olson, P. (2025): *Világuralom*. Booklab, Budapest.

Our World in Data (2026): *Global GDP over the long run*. <https://ourworldindata.org/grapher/global-gdp-over-the-long-run>. (Letöltve: 2026. feb. 1.)

Perez, C. (2002): *Technological Revolutions and Financial Capital*. Edward Elgar, Cheltenham.

Perkins, A.; Perkins, M. (2001): *The Internet Bubble*. Harperbusiness, New York.

Schwartz, P. ; Leyden, P.; Hyatt, J. (1999): *The Long Boom*. Basic Books, New York.

Solow, R. (1987): *Manufacturing Matters*. The New York Times, július 12.

Stanford University (2026): *AI Index Report 2025*. Stanford HAI, <https://aiindex.stanford.edu/report/>. (Letöltve: 2026. feb. 2.).

Suleyman, M. (2023): *The Coming Wave*. Crown, New York.

Thurow, L. (1986): *White-Collar Overhead*. Across the Board, 11.

U.S. Bureau of Labor Statistics (2026): *Productivity*. <https://www.bls.gov/productivity/>. (Letöltve: 2026. jan. 25.).

Wheeler, K. (2025): *MIT: Why 95% of Enterprise AI Investments Fail to Deliver*. AI.magazine, szeptember 8. <https://aimagazine.com/news/mit-why-95-of-enterprise-ai-investments-fail-to-deliver>. (Letöltve: 2026. jan. 25.)

World Bank (2025): *Digital Progress and Trends Report*. World Bank Group. <https://www.worldbank.org/en/publication/dptr2025-ai-foundations>. (Letöltve: 2026. jan. 30.)

World Bank (2026a): *GDP growth*. <https://data.worldbank.org/indicator/NY.GDP.MKTP.KD.ZG>. (Letöltve: 2026. feb. 2.)

World Bank (2026b): *Global Economic Prospects*. January

A tudatos mesterséges intelligencia kora: ember és gép közös döntéseinek szabályozási és gazdasági kérdései

The era of Trustworthy AI: Regulatory and economic dimensions of human-machine decision-making

Halász Rita^{1,2} 

Absztrakt

A mesterséges intelligencia (Artificial Intelligence, AI) fejlődése az utóbbi években jelentősen átalakította a gazdasági, társadalmi és döntéshozatali rendszereket. Az AI rendszerek nemcsak automatizálják a folyamatokat, hanem egyre inkább részt vesznek komplex döntéshozatali mechanizmusokban is. A tanulmány célja annak bemutatása, hogy miként alakul az ember és a mesterséges intelligencia együttműködése a döntéshozatalban, milyen alapelvek mentén értelmezhető a megbízható mesterséges intelligencia fogalma, valamint hogyan szabályozza ezt a területet az Európai Unió és Magyarország jogrendszere. A kutatás külön figyelmet fordít az AI Act kockázatalapú szabályozási modelljére, a magyar AI-szabályozás főbb elemeire, valamint a vállalati AI-alkalmazások fejlődési szintjeire.

Kulcsszavak: mesterséges intelligencia, AI Act, human-in-the-loop, digitális gazdaság, technológiai szabályozás

Abstract

The development of Artificial Intelligence (AI) in recent years has fundamentally transformed economic, social, and decision-making systems. AI systems not only automate processes but increasingly participate in complex decision-making mechanisms as well. The aim of this study is to examine how cooperation between humans and artificial intelligence is evolving in decision-making processes, to explore the principles underlying the concept of trustworthy artificial intelligence, and to analyse how this field is regulated within the legal frameworks of the European Union and Hungary. The research pays particular attention to the risk-based regulatory model of the AI Act, the main elements of Hungarian AI regulation, and the levels of development of corporate AI applications.

¹

² dr. Halász Rita jogász-közgazdász, a Blockchain Magyarország Egyesület alelnöke, a DigitalTech EDIH projektmenedzsere, valamint üzletfejlesztési, jogi compliance és stratégiai tanácsadója. Elérhetőség: rita.halasz.dr@blockchainhungary.org LinkedIn: [https://www.linkedin.com/in/rita-hal%C3%A1sz-](https://www.linkedin.com/in/rita-hal%C3%A1sz-47699041?utm_source=share_via&utm_content=profile&utm_medium=member_ios)

[47699041?utm_source=share_via&utm_content=profile&utm_medium=member_ios](https://www.linkedin.com/in/rita-hal%C3%A1sz-47699041?utm_source=share_via&utm_content=profile&utm_medium=member_ios)

Keywords: artificial intelligence, AI Act, human-in-the-loop, digital economy, technology regulation

1. Bevezetés

A mesterséges intelligencia fejlődése az információs társadalom egyik legjelentősebb technológiai fordulópontja. A digitális rendszerek fejlődése következtében az adatok mennyisége és feldolgozási képessége olyan szintet ért el, amely lehetővé teszi komplex prediktív modellek alkalmazását gazdasági és társadalmi folyamatokban. Yuval Noah Harari – izraeli történész és író – szerint az emberiség történetében először hozott létre olyan technológiát, amely bizonyos értelemben „képes jobban megérteni bennünket, mint mi saját magunkat” (Harari, 2018). Ez a megállapítás rávilágít arra, hogy az AI fejlődése nem pusztán technológiai kérdés, hanem alapvetően érinti az emberi döntéshozatal természetét is. A modern algoritmusok képesek hatalmas adatmennyiségek gyors feldolgozására, összetett mintázatok felismerésére és olyan előrejelzések készítésére, amelyek egyre inkább befolyásolják a gazdasági, üzleti és társadalmi döntéshozatalt. Ennek következtében a mesterséges intelligencia nem csupán az automatizáció eszközeként jelenik meg, hanem a döntéstámogatás és a stratégiai elemzés meghatározó technológiájává válik. A technológiai fejlődés így egy új döntéshozatali paradigma kialakulásához vezet, amelyben az emberi és a mesterséges intelligencia együttműködése kerül a középpontba. Ez a modell – amelyet a szakirodalom gyakran „human-in-the-loop” vagy „augmented intelligence” megközelítésként ír le – azt feltételezi, hogy a jövő gazdasági és társadalmi rendszereiben az ember és a gép közös döntéshozatala válik meghatározóvá.

2. A mesterséges intelligencia gazdasági és társadalmi jelentősége

A mesterséges intelligencia fejlődése napjainkra a digitális gazdaság egyik legmeghatározóbb technológiai tényezőjévé vált. A technológia jelentősége abban rejlik, hogy képes hatalmas mennyiségű adat feldolgozására, komplex mintázatok felismerésére és olyan prediktív modellek létrehozására, amelyek támogatják a gazdasági és társadalmi döntéshozatalt. Az adatokon alapuló elemzések lehetővé teszik a vállalatok számára, hogy pontosabban előre jelezzék a piaci folyamatokat, optimalizálják erőforrásaikat és hatékonyabban reagáljanak a gazdasági környezet változásaira.

Gazdasági szempontból a mesterséges intelligencia a termelékenység növekedésének egyik legfontosabb motorja lehet. A különböző gazdasági elemzések szerint az AI alkalmazása jelentős mértékben hozzájárulhat a vállalatok hatékonyságának növekedéséhez, az automatizált folyamatok elterjedéséhez és az innováció felgyorsulásához. Az algoritmikus rendszerek képesek azonosítani olyan mintázatokat és összefüggéseket, amelyek az emberi elemzők számára gyakran rejtve maradnak, ezáltal hozzájárulva a gazdasági döntések megalapozottságához. A gazdasági hatások mellett a mesterséges intelligencia jelentős társadalmi következményekkel is jár, hiszen az algoritmikus rendszerek egyre inkább jelen vannak az emberek mindennapi életében.

A mesterséges intelligencia társadalmi és gazdasági alkalmazásának egyik kulcskérdése a jövőben a hiteles adatok rendelkezésre állása és az algoritmikus rendszerek működésének bizonyíthatósága lesz. Az AI-modellek megbízhatósága nagymértékben függ a felhasznált adatok minőségétől, reprezentativitásától és ellenőrizhetőségétől. A pontatlan vagy torz adatokon alapuló modellek ugyanis hibás következtetésekhez és nem kívánt társadalmi hatásokhoz vezethetnek. Emiatt a mesterséges intelligencia fejlődésének következő szakaszában egyre nagyobb jelentőséget kap az adatok hitelességének biztosítása, az algoritmusok átlátható működése, valamint a rendszerek auditálhatósága és bizonyítható működése. Az úgynevezett „explainable AI” és „trustworthy AI” megközelítések éppen ezt a célt szolgálják, hogy olyan technológiai és szabályozási keretek alakuljanak ki, amelyek biztosítják, hogy az AI-rendszerek működése ellenőrizhető, dokumentálható és társadalmilag is elfogadható legyen.

A jövő digitális gazdaságában ezért kulcsfontosságú lesz, hogy a mesterséges intelligencia rendszerei megbízható adatokon alapuljanak, működésük átlátható legyen, és a technológia alkalmazása összhangban maradjon a társadalmi értékekkel.

3. Az ember és a mesterséges intelligencia együttműködésének fő paradigmája

Yuval Noah Harari a digitális korszak egyik legfontosabb kihívását abban látja, hogy a döntéshozatal egyre inkább az adatfeldolgozás képességéhez kötődik. Véleménye szerint a 21. században azok a társadalmak és szervezetek kerülnek előnybe, amelyek képesek hatékonyan integrálni az emberi és az algoritmikus intelligenciát. Harari megfogalmazásában a modern gazdaságban „a hatalom egyre inkább azok kezében összpontosul, akik képesek az adatokat feldolgozni és a döntéshozatalt algoritmusokkal támogatni”, ami alapvetően átalakítja a gazdasági és társadalmi struktúrákat. E gondolat jól mutatja a mesterséges intelligencia korszakának központi dilemmáját: miként alakul az ember szerepe egy olyan világban, ahol a döntések jelentős részét algoritmikus rendszerek készítik elő.

A mesterséges intelligencia fejlődése ugyanis nem csupán technológiai innovációt jelent, hanem a döntéshozatal természetének átalakulását is. A korábbi ipari forradalmak elsősorban az emberi fizikai munkát helyettesítették vagy támogatták gépekkel, míg a jelenlegi technológiai korszak az emberi kognitív képességek bizonyos aspektusait egészíti ki. A modern algoritmikus rendszerek képesek hatalmas mennyiségű adat feldolgozására, komplex statisztikai mintázatok felismerésére és előrejelzések készítésére, amelyek gyakran meghaladják az emberi elemzők kapacitását. Ennek következtében a gazdasági és intézményi döntéshozatal egyre inkább adatvezérelté válik.

A szakirodalomban az ember és a mesterséges intelligencia együttműködésére épülő döntéshozatali modellt gyakran human-in-the-loop vagy augmented intelligence megközelítésként írják le. Ebben a rendszerben a mesterséges intelligencia az adatelemzés és a mintázatfelismerés területén nyújt kiemelkedő teljesítményt, míg az ember feladata a kontextus értelmezése, az értékalapú döntések meghozatala és a felelősségvállalás. Az algoritmusok tehát nem önálló döntéshozóként jelennek meg, hanem az emberi döntések támogatásának eszközeként.

A kutatások egyre inkább azt mutatják, hogy a tisztán automatizált vagy kizárólag emberi döntési modellekkel szemben a hibrid rendszerek bizonyulnak hatékonyabbnak. Az MIT és más kutatóintézetek vizsgálatai például kimutatták, hogy az orvosi diagnosztikában az AI által támogatott orvosok gyakran pontosabb eredményeket érnek el, mint az algoritmus vagy az ember önálló alkalmazása esetén. Hasonló eredmények figyelhetők meg a pénzügyi kockázatelemzés, a logisztikai optimalizáció és az ipari termelés területén is.

Az ember és a mesterséges intelligencia együttműködésének paradigmája tehát nem a technológiai autonómia irányába mutat, hanem egy olyan új döntéshozatali modell felé, amelyben az algoritmusok és az emberi szakértelem egymást kiegészítve működnek. A mesterséges intelligencia ebben az értelemben nem az emberi döntéshozatal végét jelenti, hanem annak új formáját. A jövő gazdaságában egyre inkább azok a szervezetek lesznek sikeresek, amelyek képesek a technológiai eszközöket az emberi kreativitással, stratégiai gondolkodással és etikai felelősséggel integrálni.

Harari gondolatát tovább értelmezve megállapítható, hogy a mesterséges intelligencia korszakának legfontosabb kérdése nem az, hogy a gépek képesek lesznek-e az ember helyett dönteni, hanem az, hogy miként alakul az ember szerepe egy olyan világban, ahol az algoritmusok egyre nagyobb mértékben befolyásolják a döntési folyamatokat. A jövő gazdasági rendszereiben ezért a versenyképesség kulcsa az emberi és a mesterséges intelligencia hatékony együttműködésének kialakítása lesz. Ebben a modellben a mesterséges intelligencia az adatokból

származó mintázatokat és lehetőségeket tárja fel, míg az ember biztosítja a stratégiai irányt, az értékalapú döntéseket és a társadalmi felelősséget.

A tudatos mesterséges intelligencia korszakának központi gondolata tehát az, hogy a technológia nem az emberi döntéshozatal helyettesítésére, hanem annak kiterjesztésére és megerősítésére szolgál. Az ember és a gép közös intelligenciája olyan döntéshozatali rendszereket hozhat létre, amelyek egyszerre képesek hatékonyak, adatvezéreltek és társadalmilag felelősek lenni.

4. A mesterséges intelligencia gazdasági szerepe

Az adatvezérelt gazdaság kialakulásával az algoritmikus rendszerek nem csupán az információfeldolgozás eszközei, hanem egyre inkább a gazdasági döntéshozatal, az innováció és a versenyképesség meghatározó tényezői. A digitális gazdaságban az adatok feldolgozásának képessége stratégiai erőforrássá vált, amely jelentős versenyelőnyt biztosít a vállalatok és gazdaságok számára. A mesterséges intelligencia alkalmazása lehetővé teszi a nagy mennyiségű adat gyors és hatékony elemzését, amely hozzájárul a termelékenység növekedéséhez, az üzleti folyamatok optimalizálásához és új üzleti modellek kialakulásához.

A gazdasági elemzések egyértelműen rámutatnak arra, hogy a mesterséges intelligencia jelentős növekedési potenciállal rendelkezik. A PwC globális gazdasági előrejelzése szerint a mesterséges intelligencia alkalmazása 2030-ra akár 15,7 billió dollárral növelheti a világ gazdaság teljesítményét. Ez a növekedés elsősorban két fő tényezőtől származik: egyrészt az automatizáció által elért termelékenységnövekedésből, másrészt az AI-alapú innovációk által generált új termékekből és szolgáltatásokból. A McKinsey Global Institute kutatásai szintén azt mutatják, hogy az AI-technológiák különösen jelentős hatást gyakorolnak az ipari termelés, a pénzügyi szolgáltatások, a logisztika, valamint az egészségügyi rendszerek működésére.

A vállalatok számára a mesterséges intelligencia alkalmazása elsősorban a hatékonyság növelésében és a döntéshozatal javításában jelenik meg. Az algoritmikus rendszerek képesek gyorsan elemezni a piaci adatokat, előre jelezni a fogyasztói viselkedést és optimalizálni az ellátási láncokat. A digitális platformgazdaság szereplői – például az e-kereskedelmi vállalatok és a technológiai platformok – különösen intenzíven alkalmazzák a mesterséges intelligenciát az ügyfeladatok elemzésére, a személyre szabott ajánlórendszerek működtetésére és a marketingstratégiák optimalizálására. Az ilyen rendszerek jelentősen növelhetik a vállalatok bevételeit, mivel képesek pontosabban azonosítani a fogyasztói igényeket és hatékonyabban célozni a potenciális ügyfeleket.

A mesterséges intelligencia gazdasági szerepe különösen jól megfigyelhető az ipari termelés és az úgynevezett „ipar 4.0” rendszerek fejlődésében. Az intelligens gyártási rendszerek képesek valós időben monitorozni a termelési folyamatokat, előre jelezni a gépek meghibásodását, valamint optimalizálni az erőforrások felhasználását. A prediktív karbantartási rendszerek például lehetővé teszik, hogy a vállalatok a gépek működéséből származó adatok elemzésével előre azonosítsák a potenciális hibákat, ami jelentősen csökkentheti a leállásokból eredő költségeket.

A pénzügyi szektorban a mesterséges intelligencia elsősorban a kockázatelemzés, a csalásfelderítés és az algoritmikus kereskedelem területén játszik kulcsszerepet. Az AI-alapú rendszerek képesek hatalmas mennyiségű tranzakciós adat elemzésére, ami lehetővé teszi a pénzügyi intézmények számára, hogy gyorsabban azonosítsák a gyanús tranzakciókat és csökkentsék a csalásból eredő veszteségeket. Emellett a hitelbírálati rendszerek is egyre inkább algoritmikus modellekre épülnek, amelyek a hagyományos pénzügyi mutatók mellett számos további adatforrást is figyelembe vesznek.

Az egészségügyi szektorban a mesterséges intelligencia különösen a diagnosztikai rendszerek és a gyógyszerkutatás területén hoz jelentős változásokat. Az AI-alapú képfeldolgozó rendszerek például képesek nagy pontossággal azonosítani bizonyos betegségek jeleit orvosi képalkotó vizsgálatok során. A gyógyszerkutatásban pedig az algoritmusok felgyorsíthatják az új gyógyszermolekulák azonosítását és a klinikai kutatások tervezését.

A mesterséges intelligencia gazdasági hatása ugyanakkor nem csupán a termelékenység növekedésében jelenik meg, hanem a munkaerőpiac átalakulásában is. Az automatizáció következtében bizonyos munkakörök iránt csökkenhet a kereslet, miközben új szakmák és kompetenciák jelennek meg. A Világgazdasági Fórum előrejelzései szerint a mesterséges intelligencia és az automatizáció a következő években jelentős strukturális változásokat idézhet elő a munkaerőpiacon, ami új készségek és kompetenciák iránti igényt teremt.

A mesterséges intelligencia gazdasági szerepének egyik legfontosabb sajátossága az, hogy a technológia hálózati hatások révén erősíti a piaci koncentrációt. Azok a vállalatok, amelyek nagy mennyiségű adat felett rendelkeznek és képesek fejlett algoritmikus rendszerek fejlesztésére, jelentős versenyelőnyre tehetnek szert. Ez a jelenség különösen a digitális platformgazdaságban figyelhető meg, ahol a vezető technológiai vállalatok domináns piaci pozíciót alakítottak ki.

A technológia tehát nem csupán az üzleti folyamatok hatékonyságát növeli, hanem új gazdasági struktúrák és üzleti modellek kialakulását is elősegíti. A jövő gazdaságában ezért egyre inkább azok a vállalatok és gazdaságok lesznek sikeresek, amelyek képesek hatékonyan integrálni a mesterséges intelligenciát működésükbe, miközben megfelelően kezelik a technológia által felvetett társadalmi és szabályozási kihívásokat.

5. A vállalati AI-alkalmazások fejlődési szintjei

A mesterséges intelligencia vállalati alkalmazása az elmúlt évtizedben fokozatos fejlődési folyamaton ment keresztül. A technológia kezdetben elsősorban az adatfeldolgozás és az automatizáció eszközeként jelent meg, azonban a digitális infrastruktúra fejlődésével és a nagy mennyiségű adat rendelkezésre állásával az AI egyre inkább stratégiai jelentőségű eszközzé vált a vállalatok működésében. A vállalati AI-megoldások fejlődése több szintre bontható, amelyek jól tükrözik azt, hogyan változik a technológia szerepe az üzleti folyamatokban: az egyszerű automatizációtól a komplex döntéstámogatáson át egészen a félautonóm rendszerekig.

A fejlődés első szintjén a mesterséges intelligencia elsősorban az ismétlődő, szabályalapú feladatok automatizálására szolgál. Ebben a fázisban az AI-rendszerek célja a hatékonyság növelése és az adminisztratív terhek csökkentése. A vállalatok gyakran alkalmaznak ilyen technológiákat ügyfélszolgálati chatbotok, dokumentumfeldolgozó rendszerek vagy egyszerű adatfeldolgozó algoritmusok formájában. Ezek a rendszerek képesek gyorsan feldolgozni nagy mennyiségű információt, csökkenteni az emberi hibák lehetőségét és jelentősen mérsékelni az operatív költségeket. A digitális ügyfélkiszolgálás területén például az AI-alapú chatbotok képesek automatikusan válaszolni a gyakran ismétlődő kérdésekre, ami lehetővé teszi az emberi munkatársak számára, hogy a komplexebb ügyekre koncentráljanak.

A fejlődés második szintjén jelennek meg a döntéstámogató rendszerek, amelyek már nem csupán automatizálják az üzleti folyamatokat, hanem aktívan hozzájárulnak a vezetői döntések előkészítéséhez. Az ilyen rendszerek gyakran fejlett adat- és prediktív elemzési technológiákra épülnek, amelyek képesek a múltbeli adatokból mintázatokat felismerni és jövőbeli trendeket előre jelezni. A pénzügyi szektorban például az AI-alapú rendszerek képesek modellezni a piaci kockázatokat, előre jelezni a hitelkockázatot vagy azonosítani a potenciális csalási mintázatokat. Hasonló módon a marketing területén az algoritmusok elemezhetik a fogyasztói viselkedést, és

személyre szabott ajánlásokat generálhatnak, amelyek növelik a vásárlói elköteleződést és a bevételeket.

A döntéstámogató rendszerek egyik legfontosabb sajátossága, hogy a mesterséges intelligencia és az emberi szakértelem szoros együttműködésére épülnek. Az algoritmusok gyors és komplex adatelemzést végeznek, míg az emberi döntéshozók felelősek a stratégiai értelmezésért és a végső döntések meghozataláért. Ebben a modellben az AI a döntéshozatal hatékonyságát és megalapozottságát növeli, miközben az emberi kontroll továbbra is megmarad.

A fejlődés harmadik szintjén jelennek meg a félautonóm rendszerek, amelyek már képesek komplex üzleti folyamatok önálló elemzésére és optimalizálására, valamint bizonyos döntések automatizált meghozatalára. Ezek a rendszerek valós idejű adatfeldolgozásra és folyamatosan tanuló algoritmusokra épülnek, amelyek a működés során keletkező adatok alapján folyamatosan fejlesztik működésüket. Az ipari termelés területén például az AI-alapú rendszerek képesek optimalizálni a gyártási folyamatokat, előre jelezni a gépek meghibásodását és automatikusan módosítani a termelési paramétereket a hatékonyság növelése érdekében. A logisztikai rendszerekben az algoritmusok optimalizálhatják a szállítási útvonalakat és a készletgazdálkodást, ami jelentős költségmegtakarítást eredményezhet.

A legfejlettebb vállalati AI-rendszerek esetében a technológia már integrált módon működik a vállalat digitális infrastruktúrájában, és képes különböző üzleti folyamatok összehangolt optimalizálására. Az ilyen rendszerek gyakran több adatforrásból származó információkat integrálnak, és komplex modellek segítségével támogatják a stratégiai döntéshozatalt. A generatív mesterséges intelligencia megjelenése tovább gyorsította ezt a fejlődési folyamatot, mivel az ilyen rendszerek nemcsak adatokat elemeznek, hanem tartalmakat generálnak, kódot írnak, vagy akár üzleti stratégiák kidolgozásában is részt vehetnek.

A vállalati AI-alkalmazások fejlődési szintjei tehát jól tükrözik a technológia gazdasági jelentőségének növekedését. Míg a korai alkalmazások elsősorban a költségek csökkentésére és az operatív hatékonyság növelésére irányultak, a fejlettebb rendszerek már stratégiai értéket teremtenek a vállalatok számára. A mesterséges intelligencia integrációja ezért egyre inkább a vállalati versenyképesség egyik kulcstényezőjévé válik, különösen azokban az iparágakban, ahol a gyors döntéshozatal és az adatvezérelt működés meghatározó szerepet játszik.

A jövőben várhatóan tovább erősödik az a tendencia, hogy a mesterséges intelligencia a vállalati működés szerves részévé válik. Azok a vállalatok, amelyek képesek hatékonyan integrálni az AI-technológiákat működésükbe, jelentős versenyelőnyre tehetnek szert. Ugyanakkor a technológia alkalmazása új kihívásokat is felvet, különösen a döntéshozatal átláthatósága, az adatvédelem és az etikai felelősség kérdésében. Emiatt a vállalati AI-rendszerek fejlődése szorosan összekapcsolódik a megfelelő szabályozási és irányítási keretek kialakításával, amelyek biztosítják, hogy a technológia alkalmazása összhangban legyen a társadalmi és gazdasági érdekekkel.

6. A mesterséges intelligencia alkalmazásának kockázatai

A mesterséges intelligencia gazdasági és társadalmi alkalmazása jelentős előnyökkel jár, ugyanakkor számos olyan kockázatot is hordoz, amelyek megfelelő szabályozási és intézményi válaszokat igényelnek. Az AI-alapú rendszerek egyre szélesebb körű alkalmazása, a pénzügyi szolgáltatásoktól kezdve az egészségügyön és az ipari termelésen át egészen a közszolgáltatásokig, új típusú kockázatokhoz hoz létre, amelyek nemcsak technológiai, hanem gazdasági, jogi és etikai dimenzióval is rendelkeznek. A mesterséges intelligencia által generált döntések egyre nagyobb hatással vannak az egyének életére és a gazdasági folyamatokra, ezért

különösen fontos megérteni azokat a potenciális veszélyeket, amelyek az algoritmikus rendszerek működéséből fakadhatnak.

Az egyik leggyakrabban említett kockázat az úgynevezett algoritmikus torzítás, amely akkor jelentkezik, amikor az AI-rendszerek a tanulási adatokban meglévő torzulásokat vagy társadalmi egyenlőtlenségeket reprodukálják. Mivel a gépi tanulási modellek működése nagymértékben a rendelkezésre álló adatok minőségétől függ, az algoritmusok gyakran tükrözik az adatforrásokban meglévő torzításokat. Ez különösen problematikus lehet olyan területeken, mint a hitelbírálat, a munkaerő-kiválasztás vagy a bűnüldözés, ahol az algoritmikus döntések közvetlen hatással vannak az egyének lehetőségeire. Számos kutatás kimutatta, hogy bizonyos algoritmikus rendszerek például etnikai vagy társadalmi háttér alapján torz eredményeket generálhatnak, ami komoly etikai és jogi kérdéseket vet fel.

A mesterséges intelligencia alkalmazásával kapcsolatos másik fontos kockázat az átláthatóság hiánya, amelyet gyakran „black box” problémaként említenek. A modern mélytanulási rendszerek működése sok esetben rendkívül komplex, és még a fejlesztők számára is nehéz pontosan megmagyarázni, hogy egy adott algoritmus milyen logika alapján jutott egy adott döntésre. Ez különösen problémás lehet olyan területeken, ahol a döntések jelentős következményekkel járnak, például a pénzügyi szolgáltatások, az egészségügyi diagnosztika vagy a közszféra döntéshozatali rendszereiben. Az átláthatóság hiánya nemcsak a felhasználók bizalmát csökkentheti, hanem jogi szempontból is problémát jelenthet, hiszen a döntések indokolhatósága és ellenőrizhetősége alapvető követelmény a szabályozási rendszer(ek)ben.

A kockázatok között kiemelt szerepet kap az adatvédelem és a személyes adatok kezelése is. A mesterséges intelligencia működésének egyik alapfeltétele a nagy mennyiségű adat rendelkezésre állása, amely gyakran személyes vagy érzékeny információkat tartalmaz. Az ilyen adatok gyűjtése és feldolgozása komoly adatvédelmi kérdéseket vet fel, különösen az Európai Unió adatvédelmi szabályozási keretrendszerében. Az algoritmikus rendszerek által végzett profilalkotás vagy viselkedéselemzés például jelentős hatással lehet az egyének magánszférájára, és új típusú adatvédelmi kockázatokat teremthet.

A mesterséges intelligencia alkalmazása emellett biztonsági kockázatokat is hordozhat. Az AI-rendszerek sérülékenyek lehetnek bizonyos támadási formákkal szemben, amelyek során a támadók manipulált adatok segítségével félrevezethetik az algoritmusokat. Az ilyen támadások különösen veszélyesek lehetnek olyan rendszerek esetében, amelyek kritikus infrastruktúrák működésében játszanak szerepet, például autonóm járművek, energetikai rendszerek vagy egészségügyi technológiák esetében.

A gazdasági szempontból releváns kockázatok közé tartozik a munkaerőpiaci átalakulás is. Az automatizáció és az AI-alapú rendszerek alkalmazása számos iparágban csökkentheti az alacsonyabb képzettséget igénylő munkakörök iránti keresletet, miközben növeli a magas szintű digitális kompetenciákkal rendelkező szakemberek iránti igényt. Ez strukturális változásokat eredményezhet a munkaerőpiacon, amelyek hosszú távon hatással lehetnek a társadalmi egyenlőtlenségekre és a gazdasági stabilitásra.

További jelentős kockázatot jelent a piaci koncentráció erősödése a digitális gazdaságban. Azok a vállalatok, amelyek nagy mennyiségű adat felett rendelkeznek és fejlett AI-rendszereket képesek fejleszteni, jelentős versenyelőnyre tehetnek szert a piacon. Ez különösen a digitális platformgazdaságban figyelhető meg, ahol néhány globális technológiai vállalat domináns szerepet tölt be. Az adatokhoz és a technológiai infrastruktúrához való hozzáférés egyenlőtlen eloszlása ezért új versenypolitikai és szabályozási kérdéseket vet fel.

A mesterséges intelligencia kockázatainak kezelése ezért nem csupán technológiai, hanem szabályozási és intézményi kihívást is jelent. A jövő egyik legfontosabb kérdése ezért az lesz, hogy

miként lehet egyensúlyt teremteni az innováció ösztönzése és a társadalmi kockázatok hatékony kezelése között.

7. Nemzetközi AI-tanulmányok és kutatási eredmények: gazdasági hatások, szabályozási trendek és az ember–AI együttműködés

A mesterséges intelligencia fejlődésének globális hatásait az elmúlt években számos nemzetközi kutatóintézet és szakpolitikai szervezet vizsgálta. Ezek közül kiemelkedő jelentőségű a Stanford University által publikált AI Index Report, az OECD mesterséges intelligenciára vonatkozó alapelvei, valamint a Massachusetts Institute of Technology (MIT) kutatásai az ember és a mesterséges intelligencia együttműködéséről. Ezek a tanulmányok nemcsak a technológiai fejlődés aktuális állapotát mutatják be, hanem a mesterséges intelligencia gazdasági, társadalmi és szabályozási hatásainak átfogó értelmezését is megfogalmazzák. A nemzetközi kutatások egyértelműen rámutatnak arra, hogy a mesterséges intelligencia nem pusztán technológiai innováció, hanem a gazdasági struktúrák és a döntéshozatali rendszerek átalakulásának egyik meghatározó tényezője.

A Stanford University által publikált AI Index Report a mesterséges intelligencia fejlődésének egyik legátfogóbb globális elemzése. A jelentés évente gyűjti össze és elemzi az AI-hoz kapcsolódó legfontosabb gazdasági, technológiai és társadalmi trendeket. A kutatás külön figyelmet fordít a mesterséges intelligencia gazdasági hatásaira, az iparági alkalmazások fejlődésére, valamint az AI-innovációhoz kapcsolódó befektetések alakulására. A jelentés szerint a mesterséges intelligencia alkalmazása az elmúlt években rendkívül gyors növekedést mutatott mind a vállalati szektorban, mind a kutatás-fejlesztési tevékenységek területén. A globális AI-befektetések volumene folyamatosan növekszik, és egyre több vállalat integrál mesterséges intelligencia alapú megoldásokat működésébe.

A Stanford AI Index egyik fontos megállapítása, hogy a mesterséges intelligencia gazdasági hatása egyre inkább a termelékenység növekedésében és az új üzleti modellek kialakulásában jelenik meg. A vállalatok a technológiát elsősorban az adatelemzés, az automatizáció, valamint a döntéstámogatás területén alkalmazzák.

A generatív mesterséges intelligencia megjelenése tovább gyorsította ezt a folyamatot, mivel az ilyen rendszerek nemcsak adatokat elemeznek, hanem tartalmakat generálnak, programkódot írnak vagy akár üzleti folyamatok optimalizálásában is részt vehetnek. A jelentés ugyanakkor arra is rámutat, hogy a mesterséges intelligencia fejlődése új kockázatokat is generál, például az algoritmikus torzítás, az adatvédelem vagy a technológiai koncentráció kérdésében (megerősítve a korábban leírtakat).

A nemzetközi szabályozási gondolkodás egyik meghatározó dokumentuma az OECD Principles on Artificial Intelligence, amelyet 2019-ben fogadtak el, és amely a világ első átfogó nemzetközi AI-irányelvei közé tartozik. Az OECD-elvek célja olyan normatív keretrendszer kialakítása volt, amely elősegíti a mesterséges intelligencia felelős és emberközpontú fejlesztését. Az OECD megközelítésének középpontjában az úgynevezett human-centered AI koncepció áll, amely szerint a mesterséges intelligencia fejlesztésének és alkalmazásának az emberi jólét és a társadalmi értékteremtés szolgálatában kell állnia.

Az OECD dokumentuma több alapelvet határoz meg a mesterséges intelligencia felelős alkalmazására vonatkozóan. Ezek közé tartozik az átláthatóság biztosítása, az algoritmikus döntések elszámoltathatósága, valamint a technológia biztonságos és megbízható működése. Az OECD hangsúlyozza továbbá, hogy a mesterséges intelligencia fejlesztésének támogatnia kell a gazdasági növekedést és az innovációt, miközben minimalizálja a társadalmi kockázatokat. Ezek

az elvek jelentős hatást gyakoroltak az Európai Unió mesterséges intelligencia szabályozási megközelítésére, különösen az AI Act rendelet kialakításakor.

A mesterséges intelligencia gazdasági alkalmazásának egyik legfontosabb kérdése az ember és az algoritmikus rendszerek együttműködése a döntéshozatalban. Ezzel a kérdéssel foglalkoznak az MIT Sloan School of Management kutatásai is, amelyek az úgynevezett human–AI collaboration modell működését vizsgálják. A kutatások rámutatnak arra, hogy a mesterséges intelligencia alkalmazása nem feltétlenül az emberi döntéshozatal kiváltását jelenti, hanem sokkal inkább egy olyan hibrid döntéshozatali rendszer kialakulását, amelyben az emberi és az algoritmikus intelligencia egymást kiegészítve működik.

Az MIT kutatásai szerint a mesterséges intelligencia integrációja a vállalati döntéshozatalba akkor a leghatékonyabb, ha a szervezetek képesek összehangolni az algoritmikus elemzési képességeket az emberi szakértelemmel és tapasztalattal. Az algoritmusok kiemelkedő teljesítményt nyújtanak a nagy mennyiségű adat feldolgozásában és a mintázatok felismerésében, míg az emberi döntéshozók erőssége a kontextus értelmezése, a kreatív problémamegoldás és az etikai megfontolások figyelembevétele. A kutatások azt mutatják, hogy a human–AI együttműködés különösen hatékony lehet olyan komplex döntési helyzetekben, ahol a gyors adatfeldolgozás és a stratégiai értelmezés egyaránt szükséges.

A nemzetközi kutatások eredményei tehát megerősítik azt a megállapítást, hogy a mesterséges intelligencia fejlődése alapvetően átalakítja a gazdasági rendszereket és a döntéshozatal mechanizmusait. A Stanford AI Index a technológiai fejlődés gazdasági hatásait és innovációs trendjeit mutatja be, az OECD elvei normatív keretet biztosítanak a felelős AI-fejlesztéshez, míg az MIT kutatásai az ember és a mesterséges intelligencia együttműködésének gyakorlati működését vizsgálják. Ezek a kutatások együttesen azt a képet rajzolják ki, hogy a mesterséges intelligencia jövője nem a teljes automatizáció irányába mutat, hanem egy olyan új gazdasági és döntéshozatali modell felé, amelyben az ember és a gép közös intelligenciája válik meghatározóvá.

Ez a megközelítés összhangban áll a mesterséges intelligencia szabályozásának európai filozófiájával is, amely az innováció ösztönzése mellett kiemelt figyelmet fordít az emberi jogok, a társadalmi értékek és a gazdasági fenntarthatóság védelmére. A nemzetközi kutatások és szakpolitikai dokumentumok ezért fontos szerepet játszanak abban, hogy a mesterséges intelligencia fejlődése olyan irányba haladjon, amely hosszú távon is elősegíti a gazdasági fejlődést és a társadalmi jólétet.

8. A mesterséges intelligencia szabályozásának normatív alapjai: a megbízható AI koncepciója

A mesterséges intelligencia európai szabályozási megközelítésének kialakulása szorosan kapcsolódik az Európai Bizottság által létrehozott High-Level Expert Group on Artificial Intelligence (AI HLEG) munkájához. A szakértői csoport 2019-ben publikálta az „Ethics Guidelines for Trustworthy AI” című munkadokumentumot, amely a mesterséges intelligencia felelős fejlesztésének és alkalmazásának normatív alapjait határozza meg. A dokumentum alapfelvetése, hogy a mesterséges intelligencia rendszereknek három alapvető követelménynek kell megfelelniük: jogszerűnek, etikusnak és technikailag megbízhatónak kell lenniük. A megbízható mesterséges intelligencia tehát nem pusztán technológiai kérdés, hanem egy olyan komplex szabályozási és etikai keretrendszer, amely a technológiai innováció és a társadalmi értékek közötti egyensúlyt kívánja biztosítani.

A szakértői csoport által kidolgozott iránymutatás hét kulcsfontosságú alapelvet határoz meg, amelyek a megbízható mesterséges intelligencia működésének alapját képezik. Ezek az elvek a következők: az emberi felügyelet és kontroll biztosítása, a technikai biztonság és megbízhatóság,

az adatvédelem és az adatkezelés minősége, az átláthatóság, a méltányosság és diszkriminációmentesség, a társadalmi és környezeti jólét előmozdítása, valamint az elszámoltathatóság. Az iránymutatás célja az volt, hogy egy olyan etikai keretrendszert hozzon létre, amely a mesterséges intelligencia fejlesztésének minden szakaszában, a tervezéstől a fejlesztésen át egészen a működtetésig, irányt mutat a technológiai szereplők számára.

Az első alapelv az emberi felügyelet és kontroll biztosításának követelménye, amely szerint az AI-rendszereknek olyan módon kell működniük, hogy az ember képes legyen ellenőrizni és szükség esetén korrigálni a rendszer működését. Ez az elv szorosan kapcsolódik az úgynevezett „human-in-the-loop” vagy „human oversight” megközelítéshez, amelynek célja annak biztosítása, hogy a mesterséges intelligencia ne működjön teljesen autonóm módon olyan területeken, ahol a döntések jelentős hatással lehetnek az egyének jogaira vagy életlehetőségeire.

A második alapelv a technikai robusztusság és biztonság követelménye, amely azt hangsúlyozza, hogy az AI-rendszereknek stabilan, megbízhatóan és kiszámítható módon kell működniük. A technikai robusztusság különösen fontos a kritikus infrastruktúrák, például az energetikai rendszerek, a közlekedés vagy az egészségügyi rendszerek esetében, ahol a technológiai hibák súlyos következményekkel járhatnak.

A harmadik alapelv az adatvédelem és az adatkezelés minősége, amely a mesterséges intelligencia működésének egyik legfontosabb feltételéhez kapcsolódik: az adatokhoz. Az AI-rendszerek működésének alapja a nagy mennyiségű adat feldolgozása, ezért kulcsfontosságú, hogy az adatok gyűjtése és kezelése megfeleljen a vonatkozó adatvédelmi szabályoknak, különösen az Európai Parlament és a Tanács (EU) 2016/679 rendeletének (GDPR).

A negyedik alapelv az átláthatóság, amely azt követeli meg, hogy a mesterséges intelligencia rendszerek működése, természetesen amennyire lehetséges, érthető és ellenőrizhető legyen. Az algoritmikus döntéshozatal esetében különösen fontos, hogy a felhasználók tisztában legyenek azzal, mikor kerül sor AI-alapú döntéshozatalra, valamint, hogy a rendszerek működése auditálható legyen.

Az ötödik alapelv a méltányosság és diszkriminációmentesség, amely az algoritmikus torzítások kockázatának csökkentésére irányul. A mesterséges intelligencia rendszerek fejlesztése során biztosítani kell, hogy az algoritmusok ne reprodukálják vagy erősítsék a társadalmi egyenlőtlenségeket.

A hatodik alapelv a társadalmi és környezeti jólét előmozdítása, amely hangsúlyozza, hogy a mesterséges intelligencia alkalmazásának nem csupán gazdasági, hanem társadalmi értéket is kell teremtenie. Ez az elv különösen fontos a fenntartható fejlődés és a digitális társadalom hosszú távú működése szempontjából.

A hetedik alapelv az elszámoltathatóság, amely biztosítja, hogy az AI-rendszerek működése mögött mindig legyenek egyértelműen meghatározható felelős szereplők. Az algoritmikus döntések következményeiért a technológiát fejlesztő vagy alkalmazó szervezetek viselik a felelősséget, ezért fontos, hogy a megfelelő auditálási és felügyeleti mechanizmusok biztosítottak legyenek.

Ezek az alapelvek később meghatározó szerepet játszottak az Európai Unió mesterséges intelligencia szabályozási keretrendszerének kialakításában, különösen az AI Act elfogadásakor. Az AI Act – amely az első átfogó mesterséges intelligencia szabályozás a világon – egy kockázatalapú szabályozási modellt alkalmaz, amely közvetlenül tükrözi a megbízható mesterséges intelligencia elveit.

Az emberi felügyelet elve például az AI Act-ben konkrét kötelezettséggként jelenik meg a magas kockázatú AI-rendszerek esetében. A rendelet előírja, hogy az ilyen rendszerek működése során biztosítani kell az emberi felügyeletet és beavatkozási lehetőséget. A technikai robusztusság és biztonság követelménye szintén explicit módon megjelenik a szabályozásban, amely előírja a magas kockázatú rendszerek számára a megfelelő kockázatkezelési mechanizmusokat, a rendszeres tesztelést és a működés folyamatos monitorozását.

Az átláthatóság elve az AI Act-ben több formában is érvényesül. A rendelet például előírja, hogy bizonyos AI-rendszerek esetében a felhasználókat tájékoztatni kell arról, hogy mesterséges intelligenciával lépnek kapcsolatba. Emellett a magas kockázatú rendszerek esetében részletes dokumentációs és nyilvántartási kötelezettségek is érvényesülnek.

A méltányosság és a diszkriminációmentesség elve az AI Act-ben elsősorban az adatminőségre és az algoritmikus modellek tesztelésére vonatkozó követelményekben jelenik meg. A szabályozás megköveteli, hogy a magas kockázatú AI-rendszerek tanulási adatai megfelelő minőségűek és reprezentatívak legyenek annak érdekében, hogy minimalizálják az algoritmikus torzítás kockázatát.

Az elszámoltathatóság elve szintén kulcsfontosságú szerepet kap a rendeletben, amely világosan meghatározza az AI-rendszerek fejlesztőinek, szolgáltatóinak és felhasználóinak felelősségi körét. Az AI Act bevezeti például a megfelelőségi értékelési mechanizmusokat és a szabályozási felügyelet intézményét is.

A megbízható mesterséges intelligencia elvei tehát nemcsak normatív iránymutatásként szolgálnak, hanem konkrét jogi követelmények formájában is megjelennek az AI Act rendelkezéseiben. Ez a megközelítés jól tükrözi az európai technológiai politika alapvető célját: a mesterséges intelligencia innovációjának ösztönzését úgy, hogy közben megőrzi az alapvető jogokat, az emberi méltóságot és a demokratikus értékeket.

9. Az Európai Unió és Magyarország mesterséges intelligencia szabályozása, és a kockázatalapú modell gyakorlati érvényesülése

A mesterséges intelligencia megoldások egyre szélesebb körű alkalmazása, a gazdasági döntéshozataltól kezdve a közszolgáltatásokon át egészen a mindennapi digitális platformokig, szükségessé tette olyan jogi és intézményi keretek kialakítását, amelyek egyszerre képesek támogatni a technológiai innovációt és kezelni az AI-rendszerek által generált társadalmi, gazdasági és etikai kockázatokat. Az Európai Unió erre a kihívásra válaszul átfogó szabályozási megközelítést dolgozott ki a mesterséges intelligencia területén, amelynek központi eleme az Európai Parlament és a Tanács (EU) 2024/1689 rendelete, az úgynevezett AI Act. A rendelet célja, hogy egységes jogi keretet biztosítson a mesterséges intelligencia fejlesztésére és alkalmazására az Európai Unió belső piacán, miközben garantálja az alapvető jogok, az emberi méltóság és az európai értékek védelmét.

Az AI Act egyik legfontosabb sajátossága az úgynevezett kockázatalapú szabályozási modell, amely abból a felismerésből indul ki, hogy a különböző mesterséges intelligencia alkalmazások eltérő mértékű kockázatot jelentenek a társadalom, az egyének alapvető jogai és a gazdasági rendszerek működése szempontjából. Ennek megfelelően a szabályozás nem egységes módon kezeli az összes AI-rendszert, hanem azok potenciális hatása és kockázati szintje alapján differenciált követelményeket határoz meg. A rendelet négy fő kockázati kategóriát különböztet meg: minimális kockázatú általános MI alkalmazások, korlátozott kockázatú bizonyos MI alkalmazások, nagy kockázatú és tiltott MI rendszereket.

A minimális kockázatú rendszerek közé tartoznak azok az AI-alkalmazások, amelyek használata nem jelent jelentős veszélyt az egyének jogaira vagy a társadalmi rendszerek működésére. Ilyen rendszerek például az online platformok ajánlórendszerei, a spam-szűrők vagy a digitális marketingben alkalmazott algoritmusok. Ezek esetében a szabályozás alapvetően nem ír elő külön kötelezettségeket, mivel az alkalmazásuk alacsony kockázattal jár, és a szabályozási megközelítés elsősorban az innováció szabadságának megőrzésére törekszik.

A korlátozott kockázatú rendszerek esetében már megjelennek bizonyos átláthatósági követelmények. Ezek jellemzően olyan AI-alkalmazások, amelyek közvetlen kapcsolatba lépnek a felhasználókkal, ezért fontos, hogy a felhasználók tisztában legyenek azzal, ha mesterséges intelligenciával kommunikálnak. Ebbe a kategóriába tartoznak például a chatbotok vagy a generatív mesterséges intelligencia rendszerei. A szabályozás előírja például, hogy a felhasználókat megfelelően tájékoztatni kell arról, ha egy szolgáltatás AI-alapú rendszert használ, illetve bizonyos esetekben biztosítani kell a generált tartalmak egyértelmű megjelölését is.

Az AI Act szabályozási rendszerének egyik legfontosabb eleme a magas kockázatú AI-rendszerekre vonatkozó követelményrendszer. Ezek olyan alkalmazások, amelyek jelentős hatással lehetnek az egyének jogaira, a gazdasági lehetőségekre vagy akár az emberi biztonságra. Ide tartoznak például az egészségügyi diagnosztikai rendszerek, a hitelbírálati algoritmusok, a kritikus infrastruktúrák működését támogató rendszerek vagy az autonóm járművek irányítására szolgáló technológiák. Az ilyen rendszerek esetében a szabályozás részletes követelményeket határoz meg, amelyek kiterjednek többek között a kockázatkezelési mechanizmusokra, az adatminőség biztosítására, a technikai dokumentációra, az algoritmusok átláthatóságára és az emberi felügyelet biztosítására. A fejlesztőknek és üzemeltetőknek emellett biztosítaniuk kell, hogy a rendszerek működése dokumentált, auditálható és megfelelően tesztelt legyen, valamint, hogy az algoritmusok működése ne eredményezzen diszkriminatív vagy torz döntéseket.

A kockázatalapú modell negyedik kategóriájába tartoznak az úgynevezett tiltott rendszerek, amelyek alkalmazását az Európai Unió teljes mértékben tiltja (AI Act II. Fejezete). Ezek olyan AI-technológiák, amelyek súlyosan sérthetik az alapvető jogokat vagy jelentős társadalmi kockázatot hordoznak. Ide tartoznak például a társadalmi pontozási rendszerek, amelyek az egyének viselkedése alapján értékelik a társadalmi megbízhatóságot, valamint bizonyos manipulatív algoritmusok, amelyek az emberi döntéshozatalt pszichológiai vagy viselkedési befolyásolással torzíthatják. A szabályozás emellett szigorúan korlátozza bizonyos biometrikus megfigyelési rendszerek alkalmazását is, különösen azokat, amelyek tömeges vagy valós idejű megfigyelésre alkalmasak.

Az AI Act nem önálló szabályozási keretként működik, hanem az Európai Unió szélesebb digitális szabályozási ökoszisztémájába illeszkedik. Ide tartozik többek között az általános adatvédelmi rendelet (GDPR), amely a személyes adatok kezelésének szabályait határozza meg, valamint a Digital Services Act (DSA) és a Digital Markets Act (DMA), amelyek a digitális platformok működését és a digitális piacok versenyfeltételeit szabályozzák. Ezek a jogszabályok együtt alkotják az Európai Unió digitális gazdaságának átfogó szabályozási keretrendszerét.

Magyarország mesterséges intelligenciára vonatkozó szabályozási és szakpolitikai keretrendszere szorosan illeszkedik az Európai Unió digitális politikájához. Az ország már viszonylag korán felismerte a mesterséges intelligencia gazdasági és társadalmi jelentőségét, amelynek eredményeként 2020-ban elfogadásra került a Magyarország Mesterséges Intelligencia Stratégiája (2025. szeptemberében megújítva). A stratégia a 2020–2030 közötti időszakra határozza meg az ország AI-fejlesztési irányait, és célja a mesterséges intelligencia alkalmazásának ösztönzése a gazdaság különböző szektoraiban, valamint a kutatás-fejlesztési és innovációs kapacitások erősítése.

A stratégia megvalósításának koordinálására létrejött a Mesterséges Intelligencia Koalíció, amely olyan együttműködési platformot biztosít az állami intézmények, az egyetemek, a kutatóintézetek és a vállalati szektor szereplői számára, amelynek célja a mesterséges intelligencia hazai ökoszisztémájának fejlesztése. A koalíció munkájában több száz szervezet vesz részt, és számos szakpolitikai kezdeményezést indított a mesterséges intelligencia fejlesztésének támogatására. Megjegyzendő, hogy a 2025-ben megújított stratégia és az EU AI Act rendelete alapján 2025. december 15-én megalakult a rendelet és a hazai MI stratégia végrehajtására, és a hazai ökoszisztéma szakpolitikai irányítására a Magyar Mesterséges Intelligencia Tanács (MMIT).

A magyar szabályozási környezet³ jelenleg elsősorban az európai uniós jogszabályokra épül, különösen az AI Act és a GDPR rendelkezéseire. Az AI Act rendelet közvetlenül alkalmazandó jogszabályként Magyarországon is hatályba lépett (2025. december 2-án), ami azt jelenti, hogy a mesterséges intelligencia fejlesztésével és alkalmazásával foglalkozó magyar vállalatoknak és intézményeknek is meg kell felelniük az európai szabályozásban meghatározott követelményeknek. Emellett a magyar hatóságok feladata lesz a rendelet végrehajtásának felügyelete és a megfelelés ellenőrzése.

10. Az Egyesült Államok és Kína AI-szabályozásának eltérő megközelítései

A mesterséges intelligencia szabályozása globális szinten eltérő modellek mentén alakul, amelyek mögött különböző gazdasági, politikai és technológiai stratégiák húzódnak meg. Az Európai Unió szabályozási megközelítése elsősorban az alapjogok védelmére és a megbízható mesterséges intelligencia elveire épül, míg a világ két másik meghatározó technológiai hatalma, az Egyesült Államok és Kína eltérő hangsúlyokat alkalmaz a mesterséges intelligencia szabályozásában. A két ország szabályozási modellje jelentősen különbözik egymástól mind az állami szerepvállalás mértéke, mind a technológiai fejlesztés ösztönzésének módja, mind pedig a társadalmi kontroll és az adatkezelés szabályozása tekintetében.

Az Egyesült Államok mesterséges intelligenciára vonatkozó szabályozási megközelítése alapvetően piacvezérelt és innovációközpontú. Az amerikai technológiai politika hagyományosan a vállalati innováció ösztönzésére és a technológiai versenyképesség fenntartására épül. Ennek megfelelően az Egyesült Államokban a mesterséges intelligencia szabályozása hosszú ideig elsősorban szektorális és decentralizált módon valósult meg. Az egyes iparágak – például a pénzügyi szektor, az egészségügy vagy az autonóm járművek – saját szabályozási kereteik alapján alkalmazzák az AI-technológiákat, miközben a szövetségi kormány inkább iránymutatásokat és stratégiai programokat dolgoz ki.

Az amerikai szabályozási gondolkodás egyik fontos mérföldköve a National AI Initiative Act elfogadása volt 2020-ban, amely az Egyesült Államok mesterséges intelligencia kutatásának és fejlesztésének koordinációját kívánta erősíteni. A törvény célja a kutatás-fejlesztési kapacitások bővítése, az akadémiai és ipari együttműködés támogatása, valamint az Egyesült Államok technológiai vezető szerepének megőrzése a globális AI-versenyben. Az amerikai megközelítésben a szabályozás fő célja az innováció ösztönzése és a technológiai versenyelőny fenntartása, miközben a részletes jogi keretek gyakran az iparági önszabályozásra vagy a technológiai vállalatok felelősségére épülnek.

Az Egyesült Államokban az utóbbi években ugyanakkor egyre nagyobb figyelem irányul az AI-rendszerek etikai és társadalmi hatásaira. Ennek eredményeként a Biden-adminisztráció 2023-ban kiadta az úgynevezett AI Bill of Rights iránymutatását, amely olyan alapelveket határoz meg, mint az algoritmikus diszkrimináció elleni védelem, az adatvédelem, valamint az automatizált

³ az Európai Unió mesterséges intelligenciáról szóló rendeletének magyarországi végrehajtásáról szóló 2025. évi LXXV. törvény

rendszerek átláthatósága. Ezek az iránymutatások azonban elsősorban ajánlás jellegűek, és nem alkotnak olyan átfogó, kötelező erejű szabályozási keretet, mint az Európai Unió AI Act rendelete.

Kína mesterséges intelligenciára vonatkozó szabályozási modellje ezzel szemben erősen államközpontú és stratégiai jellegű. A kínai kormány a mesterséges intelligenciát a gazdasági fejlődés és a geopolitikai verseny egyik kulcstechnológiájaként kezeli. Ennek megfelelően a kínai állam jelentős erőforrásokat fordít az AI-kutatás és fejlesztés támogatására, miközben szoros szabályozási kontrollt gyakorol a technológia alkalmazása felett.

A kínai mesterséges intelligencia stratégia egyik alapidokumentuma a New Generation Artificial Intelligence Development Plan, amelyet 2017-ben fogadtak el. A stratégia célja, hogy Kína 2030-ra a mesterséges intelligencia globális vezető hatalmává váljon. Ennek érdekében az állam jelentős beruházásokat hajt végre az AI-kutatás, az oktatás és a technológiai infrastruktúra fejlesztése területén.

A kínai szabályozási megközelítés sajátossága, hogy az állam aktívan szabályozza az AI-rendszerek működését, különösen az online platformok és a generatív mesterséges intelligencia alkalmazásai esetében. Az elmúlt években Kína több olyan szabályozást vezetett be, amely a deepfake technológiák és a generatív AI-rendszerek működésére vonatkozik. Ezek a szabályozások részletes követelményeket írnak elő az algoritmusok átláthatóságára, a tartalommoderációra és az adatkezelésre vonatkozóan.

A kínai modell ugyanakkor jelentősen eltér az európai szabályozási filozófiától, mivel a szabályozás egyik fontos célja a társadalmi stabilitás és az információs kontroll biztosítása. Az AI-technológiák alkalmazása ezért Kínában gyakran összekapcsolódik az állami felügyeleti rendszerekkel és a digitális kormányzás eszközeivel.

A három szabályozási modell jól mutatja, hogy a mesterséges intelligencia globális szabályozása eltérő értékekre és prioritásokra épül. Az Európai Unió a jogok és a társadalmi bizalom védelmét, az Egyesült Államok az innováció és a piaci verseny ösztönzését, míg Kína az állami stratégiai irányítást és a technológiai szuverenitást helyezi a középpontba.

A mesterséges intelligencia szabályozásának jövője szempontjából ezért kulcsfontosságú kérdés, hogy ezek a különböző modellek milyen mértékben közelednek egymáshoz, illetve hogyan alakul ki egy globális szabályozási környezet a mesterséges intelligencia fejlesztésére és alkalmazására. A technológia globális jellege miatt a nemzetközi együttműködés egyre fontosabbá válik annak érdekében, hogy a mesterséges intelligencia fejlődése globálisan szinten is biztonságos és fenntartható módon történjen.

11. Az ember szerepe a mesterséges intelligencia korszakában

A mesterséges intelligencia gyors fejlődése alapvetően új kérdéseket vet fel az ember szerepéről a gazdasági és társadalmi rendszerekben. A digitális technológiák korábbi hullámai elsősorban az emberi fizikai munkát támogatták vagy váltották ki, azonban a mesterséges intelligencia megjelenése már az emberi kognitív tevékenységek bizonyos aspektusait is érinti. Az algoritmikus rendszerek képesek komplex adatmintázatok felismerésére, prediktív elemzések készítésére és egyre inkább részt vesznek a döntéshozatali folyamatokban. Ez a fejlődés szükségessé teszi annak újragondolását, hogy milyen szerepet töltsön be az ember egy olyan gazdasági környezetben, ahol a döntések jelentős részét algoritmikus rendszerek támogatják.

A mesterséges intelligencia korszakában az ember szerepe nem szűnik meg, hanem átalakul. A modern gazdaságban egyre inkább egy hibrid döntéshozatali modell alakul ki, amelyben az emberi és a mesterséges intelligencia együttműködése válik meghatározóvá. Az algoritmusok hatékonyan képesek nagy mennyiségű adat feldolgozására és statisztikai mintázatok

felismerésére, ugyanakkor az ember továbbra is kulcsszerepet játszik a kontextus értelmezésében, a stratégiai gondolkodásban és az értékalapú döntések meghozatalában. Az emberi döntéshozók képesek olyan komplex társadalmi, gazdasági és etikai szempontokat figyelembe venni, amelyek gyakran túlmutatnak az algoritmikus modellek által feldolgozott adatokon.

Az ember szerepének egyik legfontosabb dimenziója a stratégiai irányítás és felelősségvállalás. A mesterséges intelligencia rendszerek működése mindig emberi tervezés és fejlesztés eredménye, ezért a technológia alkalmazásának következményeiért végső soron az emberek (fejlesztők, vállalatok és intézmények) viselik a felelősséget. Ez különösen fontos olyan területeken, ahol az algoritmikus döntések jelentős hatással lehetnek az egyének életére, például a pénzügyi szolgáltatások, az egészségügy vagy a közszolgáltatások területén. Az emberi felügyelet biztosítása ezért a mesterséges intelligencia szabályozásának egyik alapvető követelménye, amely az Európai Unió AI Act rendeletében is megjelenik.

Az ember szerepe a mesterséges intelligencia korszakában szorosan kapcsolódik az etikai döntéshozatal kérdéséhez is. Az algoritmusok működése statisztikai modelleken és adatokon alapul, ezért gyakran nem képesek figyelembe venni azokat az erkölcsi vagy társadalmi szempontokat, amelyek bizonyos döntések esetében meghatározó jelentőségűek lehetnek. Az etikai dilemmák, például az autonóm rendszerek döntései, az algoritmikus diszkrimináció vagy a technológia társadalmi hatásai, olyan kérdéseket vetnek fel, amelyek kezelése továbbra is emberi felelősség marad. Az emberi döntéshozók feladata ezért nemcsak a technológia alkalmazása, hanem annak biztosítása is, hogy a mesterséges intelligencia működése összhangban legyen a társadalmi értékekkel és az alapvető jogokkal.

A mesterséges intelligencia korszakában az ember szerepe a kreativitás és az innováció területén is kiemelt jelentőségű marad. Az algoritmusok képesek hatékonyan feldolgozni az adatokat és optimalizálni a folyamatokat, azonban az új ötletek generálása, a komplex problémák kreatív megoldása és az új stratégiai irányok meghatározása továbbra is elsősorban az emberi gondolkodáshoz kapcsolódik. A mesterséges intelligencia ezért sok esetben nem helyettesíti, hanem kiegészíti az emberi kreativitást, például az innovációs folyamatok támogatásával vagy az információk rendszerezésével.

A munkaerőpiac szempontjából a mesterséges intelligencia fejlődése jelentős strukturális változásokat hoz. Az automatizáció bizonyos rutinszerű feladatok iránt csökkentheti a keresletet, ugyanakkor új szakmák és kompetenciák iránti igényt is teremt. A digitális gazdaságban egyre fontosabbá válik az úgynevezett AI-műveltség, vagyis az a képesség, hogy az egyének képesek legyenek megérteni és hatékonyan alkalmazni a mesterséges intelligencia rendszereket. Az oktatás és a képzési rendszerek ezért kulcsszerepet játszanak abban, hogy a társadalom alkalmazkodni tudjon a technológiai változásokhoz.

Az ember szerepének átalakulása a vezetői döntéshozatal területén is megfigyelhető. A vállalati vezetők egyre gyakrabban támaszkodnak mesterséges intelligencia alapú döntéstámogató rendszerekre, amelyek képesek komplex gazdasági és piaci adatok elemzésére. A vezetői szerep azonban ebben a környezetben nem csökken, hanem inkább átalakul. A vezetők feladata egyre inkább az algoritmikus elemzések értelmezése, a stratégiai döntések meghozatala és az AI-rendszerek felelős alkalmazásának biztosítása.

A mesterséges intelligencia korszakában tehát az ember szerepe átalakul és új dimenziókat nyer. A technológia fejlődése egy olyan gazdasági és társadalmi környezetet hoz létre, amelyben az emberi és a mesterséges intelligencia együttműködése válik meghatározóvá. Az ember továbbra is kulcsszerepet játszik a stratégiai döntéshozatalban, az etikai iránymutatásban és a technológia társadalmi hatásainak kezelésében. Ebben az értelemben a mesterséges intelligencia nem az

ember helyettesítését, hanem az emberi képességek kiterjesztését jelenti, amely új lehetőségeket teremt a gazdasági fejlődés és az innováció számára.

12. Következtetések

A mesterséges intelligencia fejlődése a 21. század egyik legmeghatározóbb technológiai és gazdasági átalakulását indította el. Az algoritmikus rendszerek egyre nagyobb szerepet játszanak a gazdasági folyamatokban, a vállalati döntéshozatalban és a társadalmi rendszerek működésében. A digitális gazdaság fejlődése következtében az adatok feldolgozása és az algoritmikus elemzés képessége stratégiai jelentőségű erőforrássá vált, amely jelentős versenyelőnyt biztosíthat a vállalatok és gazdaságok számára. A mesterséges intelligencia ezért nem csupán technológiai innovációként értelmezhető, hanem a gazdasági és társadalmi rendszerek egyik alapvető strukturális átalakító tényezőjeként.

A tanulmány bemutatta, hogy a mesterséges intelligencia fejlődése új döntéshozatali modellt hoz létre, amelyben az ember és az algoritmikus rendszerek együttműködése válik meghatározóvá. A modern gazdaságban egyre inkább olyan hibrid döntéshozatali rendszerek jelennek meg, amelyek az emberi szakértelmet és a mesterséges intelligencia adatfeldolgozási képességeit kombinálják. Az algoritmusok képesek gyorsan feldolgozni nagy mennyiségű adatot és felismerni az összetett mintázatokat, míg az ember továbbra is kulcsszerepet játszik a stratégiai gondolkodásban, az értékalapú döntésekben és az etikai felelősségvállalásban. A mesterséges intelligencia tehát nem az emberi döntéshozatal megszűnését jelenti, hanem annak átalakulását és kiterjesztését.

A tanulmány rámutatott arra is, hogy a mesterséges intelligencia gazdasági hatása rendkívül jelentős. A technológia alkalmazása hozzájárulhat a termelékenység növekedéséhez, az üzleti folyamatok optimalizálásához és új gazdasági modellek kialakulásához. A vállalatok egyre szélesebb körben alkalmaznak AI-alapú rendszereket az adatelemzés, a logisztikai optimalizáció, a pénzügyi kockázatelemzés vagy a marketing területén. Ugyanakkor a mesterséges intelligencia fejlődése új kockázatokot is generál, például az algoritmikus torzítás, az adatvédelem, a technológiai koncentráció vagy a munkaerőpiac átalakulása területén.

A technológia fejlődése ezért szükségessé tette olyan szabályozási keretek kialakítását, amelyek biztosítják a mesterséges intelligencia felelős alkalmazását. A tanulmány bemutatta, hogy az Európai Unió mesterséges intelligencia szabályozási modellje a megbízható mesterséges intelligencia koncepciójára épül, amely az emberi felügyelet, az átláthatóság, az adatvédelem és az elszámoltathatóság alapelveit hangsúlyozza. Az AI Act rendelet egy kockázatalapú szabályozási modellt vezet be, amely differenciált módon kezeli a különböző AI-alkalmazásokat, és különös figyelmet fordít azokra a rendszerekre, amelyek jelentős hatással lehetnek az egyének jogaira vagy a társadalmi rendszerek működésére.

A nemzetközi kutatások – például a Stanford AI Index, az OECD mesterséges intelligenciára vonatkozó alapelvei vagy az MIT human-AI együttműködésről szóló kutatásai – szintén azt mutatják, hogy a mesterséges intelligencia jövője nem a teljes automatizáció irányába mutat, hanem egy olyan gazdasági és társadalmi modell felé, amelyben az ember és a gép együttműködése válik meghatározóvá. A mesterséges intelligencia fejlődése ezért nem csupán technológiai, hanem társadalmi és intézményi kihívás is, amely újfajta gondolkodást igényel a szabályozás, az oktatás és a gazdasági stratégiák területén. A jövő gazdaságának és társadalmának sikeressége nagymértékben attól függ majd, hogy a technológiai fejlődést milyen mértékben sikerül felelős szabályozási keretekkel és tudatos társadalmi alkalmazással összekapcsolni.

Irodalomjegyzék

Brynolfsson, E. – McAfee, A. (2014): *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company. New York.

European Commission (2019): *Ethics Guidelines for Trustworthy AI*. High-Level Expert Group on Artificial Intelligence. Brussels. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

European Commission (2020): *White Paper on Artificial Intelligence: A European Approach to Excellence and Trust*. Brussels.

European Parliament – Council of the European Union (2016): Regulation (EU) 2016/679 of the European Parliament and of the Council (General Data Protection Regulation – GDPR). Brussels.

European Parliament – Council of the European Union (2022): Regulation (EU) 2022/2065 on a Single Market for Digital Services (Digital Services Act). Brussels.

European Parliament – Council of the European Union (2022): Regulation (EU) 2022/1925 on Contestable and Fair Markets in the Digital Sector (Digital Markets Act). Brussels.

European Parliament – Council of the European Union (2024): Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act). Brussels.

Floridi, L. et al. (2018): AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. Springer Nature. Dordrecht, The Netherlands. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>

Harari, Y. N. (2016): *Homo Deus: A Brief History of Tomorrow*. Harvill Secker. London.

Harari, Y. N. (2018): *21 Lessons for the 21st Century*. Jonathan Cape. London.

McKinsey Global Institute (2023): *The Economic Potential of Generative AI: The Next Productivity Frontier*. McKinsey & Company. New York.

MIT Sloan School of Management (2023): *Expanding AI Impact with Organizational Learning*. MIT Sloan Management Review. Cambridge. <https://sloanreview.mit.edu/projects/expanding-ai-impact-with-organizational-learning/>

OECD (2019): *OECD Principles on Artificial Intelligence*. OECD Publishing. Paris. <https://oecd.ai/en/ai-principles>

OECD (2023): *Artificial Intelligence in Society*. OECD Publishing. Paris.

PwC (2017): *Sizing the Prize: What's the Real Value of AI for Your Business and How Can You Capitalise?* PricewaterhouseCoopers. London.

Russell, S. (2019): *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking. New York.

Stanford University (2024): *AI Index Report 2024*. Stanford Institute for Human-Centered Artificial Intelligence. Stanford, California. <https://aiindex.stanford.edu/report/>

World Economic Forum (2023): *The Future of Jobs Report 2023*. Geneva.

Magyarország Kormánya (2020): *Magyarország Mesterséges Intelligencia Stratégiája 2020–2030*. Budapest.

Mesterséges Intelligencia Koalíció (2021): *A mesterséges intelligencia hazai ökoszisztémájának fejlesztése*. Budapest.

Algoritmusok a csatatéren: A mesterséges intelligencia architektúráinak hatása a modern hadijogra és a döntéshozatali mechanizmusokra

Algorithms on the Battlefield: The Impact of Artificial Intelligence Architectures on the Modern Law of Armed Conflict and Decision-Making Mechanisms

Vécsey Richárd Ádám^{1,2} 

Absztrakt

A tanulmány a mesterséges intelligencia katonai alkalmazásának technológiai, döntéseméleti és jogi aspektusait vizsgálja integrált megközelítésben. A szerző elemzi a mélytanuló architektúrák szerepét a modern hadviselésben, kitérve a technológiai korlátokra, mint a domain shift vagy az adversarial támadások. A tanulmány központi tudományos újdonsága a John Boyd-féle OODA-ciklus kiterjesztése. A javasolt OODA(B) modell a visszacsatolt tanulás (backpropagation) beemelésével teszi alkalmassá a klasszikus keretrendszert a modern MI-rendszerek adaptív működésének leírására. A munka részletesen tárgyalja az algoritmikus döntéshozatal hatását a nemzetközi hadijogra, különös tekintettel az eszkalációs kockázatokra és a felelősség kérdésére. A tanulmány öt pillér, a predikció, a prevenció, az eszkaláció, a katalizáció és a redempció mentén vizsgálja fel a jövő hadviselésének stratégiai kereteit és összegzi a felmerülő problémákat.

Kulcsszavak: mesterséges intelligencia, algoritmikus döntéshozatal, autonóm fegyverrendszerek, stratégiai eszkaláció, hadijog, OODA(B) modell

Abstract

This study examines the technological, decision-theoretical, and legal aspects of the military application of artificial intelligence through an integrated approach. The author analyzes the role of deep learning architectures in modern warfare, addressing technological limitations such as domain shift and adversarial attacks. The primary theoretical contribution of the paper is the extension of John Boyd's OODA loop. The proposed OODA(B) model incorporates backpropagation, making the classical framework suitable for describing the adaptive operation of modern AI systems. The work discusses in detail the impact of algorithmic decision-making on international humanitarian law, with particular focus on escalation risks and the issue of accountability. Finally, the study outlines the strategic framework of future warfare through five pillars, such as prediction, prevention, escalation, catalyzation, and redemption, and summarizes the emerging challenges as well.

Keywords: artificial intelligence, algorithmic decision-making, autonomous weapon systems, strategic escalation, international humanitarian law, OODA(B) model

¹

² dr. Vécsey Richárd Ádám, mesterséges intelligencia fejlesztő és tanácsadó, adattudós, jogász-közigazgász
vecsey.richard@gmail.com
ORCID: (0009-0006-4203-1819) <https://orcid.org/0009-0006-4203-1819>
MTMT: (10067024) <https://m2.mtmt.hu/gui2/?type=authors&mode=browse&sel=10067024>

1. Bevezetés

A mesterséges intelligencia (MI) fejlődése alapjaiban formálja át a modern hadviselés technológiai és stratégiai környezetét. A mélytanuló architektúrák, a nagyméretű nyelvi modellek és a megerősítéses tanulás ma már a katonai döntéshozatal központi elemei. A NATO és a globális nagyhatalmak stratégiái egyaránt kiemelik az MI-t mint a következő évtized meghatározó katonai képességét, amely a felderítéstől a logisztikán át az autonóm fegyverrendszerekig minden területet átalakít.

A technológiai áttörések mögött olyan matematikai és számításelméleti alapok állnak, amelyek lehetővé teszik a komplex, dinamikus és bizonytalan környezetekben történő adaptív működést. A gradiensalapú optimalizáció, a nagy dimenziójú reprezentációs terek, a valószínűségi modellek és a többügynökös tanulási rendszerek mind hozzájárulnak ahhoz, hogy az MI képes legyen a katonai műveletekben hagyományosan emberi kompetenciának számító feladatok ellátására. A modern architektúrák olyan feldolgozási kapacitást biztosítanak, amely a csatatér információs túlterheltségét kezelhetővé teszi, és újfajta döntéshozatali mechanizmusokat szolgáltat.

A kutatás kiemelt elméleti hozzájárulása a klasszikus katonai döntéshozatali modell kiterjesztése. Mivel az MI-rendszerek alapvető sajátossága a folyamatos, adatalapú önjavítás, a tanulmány bevezeti az OODA(B) modellt. Ebben a ciklusban a hagyományos megfigyelés, orientáció, döntés és cselekvés négyesét a backpropagation visszatanítás, vagy visszacsatolt tanulás komponense egészíti ki, amely pontosabban reprezentálja a gépi tanuláson alapuló rendszerek iteratív fejlődését és adaptivitását.

Az MI alkalmazása azonban túlmutat a technológián a nemzetközi jog teljes spektrumát érintve. A klasszikus hadijogi keretrendszer négyes felosztása mentén a tanulmány elemzi a stratégiai stabilitást fenyegető jus contra bellum kockázatokat, az „algoritmikus önvédelem” jus ad bellum kérdéseit, a megkülönböztetés és arányosság jus in bello elveit, továbbá a felelősségre vonás jus post bellum dimenzióit.

2. A mesterséges intelligencia technológiai és matematikai alapjai

A modern hadviselésben alkalmazott mesterséges intelligencia olyan matematikai és architektúráis alapokra épül, amelyek lehetővé teszik a komplex, gyorsan változó és ellenséges környezetben történő adaptív működést. A katonai műveletek során keletkező adatok mennyisége, heterogenitása és időkritikussága olyan kihívásokat teremt, amelyek csak fejlett gépi tanulási módszerekkel kezelhetők. A hadszíntér sajátosságai, például a kommunikációs zavarás, a szenzorhibák, az ellenséges fél aktív megtévesztése vagy a fizikai környezet kiszámíthatatlansága tovább növelik a robusztus, decentralizált és valós idejű MI-rendszerek iránti igényt.

Az elmúlt évek hatalmas technológiai fejlődése lehetővé tette, hogy viszonylag olcsón, kis méretben, nagy számítási kapacitásokat lehessen allokálni egy-egy feladatra. A szükséges számítási műveletek elvégezhetők központi szerveren, lokális központokban, de akár eszköz-oldalon is.

A vizsgálat két szinten végezhető el. Egyrészt mikro-szinten, ahol egy-egy modell vagy döntési folyamat működése kerül elemzésre, másrészt makro-szinten, ahol a teljes hadművelleti rendszerre gyakorolt hatások értékelése történik.

2.1. A gépi tanulás matematikai alapjai

A gépi tanulás központi kérdése egy optimalizációs feladat megoldása. A modell célja, hogy a bemeneti adatok és a kívánt kimenetek közötti kapcsolatot úgy tanulja meg, hogy a

veszteségfüggvény értéke minimális legyen.³ A katonai alkalmazásokban ez a feladat sokféle formát ölthet. A célpont-azonosítás során például a modellnek minimalizálnia kell a téves felismerések számát, míg autonóm járművek esetében a navigációs hibák csökkentése a cél. Kiberhadviselési környezetben az osztályozási pontosság növelése vagy a sérülékenységek felismerésének javítása válik elsődlegessé. Szöveges elemzés és generálás esetén törekedni kell az információ pontos szintetizálására és az ellentmondó információk kezelésére, kiszűrésére.⁴

A katonai környezetben azonban a gradiensalapú tanulás számos kihívással szembesül. A szenzoradatok gyakran zajosak vagy hiányosak, a környezet változékonysága miatt domain shift lép fel, és az ellenséges fél aktív manipulációs kísérletei torzíthatják a bemeneteket. Emellett a valós idejű követelmények és a korlátozott számítási kapacitás miatt a modelleknek gyakran könnyített, optimalizált⁵ formában kell működniük.

A veszteségfüggvény megválasztása alapvetően meghatározza a modell viselkedését. A célpont-azonosítás során a cross-entropy veszteségfüggvény kiválóan alkalmazható⁶ egyedül, vagy kombinált veszteségfüggvényekkel együtt, míg autonóm járművek navigációs feladataiban a kritikák és alternatív megoldások⁷ ellenére a Huber-loss széles körben elterjedt⁸ megoldás. A katonai környezetben a veszteségfüggvények kétféle értelemben is megjelennek. Elsősorban kockázatérzékeny formában, mivel egy téves pozitív riasztás eskalációs kockázatot hordozhat, míg egy téves negatív felismerés a fenyegetés elmulasztását eredményezheti. Másodsorban pedig egy téves klasszifikáció esetén rossz célpont, például civilek vagy baráti egységek semmisülhetnek meg, illetve egy hibás döntésből fakadó navigációs probléma miatt megsemmisülhet maga az eszköz is.

A túlilleszkedés különösen veszélyes a hadszíntéren, mert a környezet folyamatosan változik. A terep, az időjárás, az ellenséges taktika vagy a szenzorok elhelyezkedése mind olyan tényezők, amelyek eltérhetnek a tanulási adatoktól. Ennek megelőzésére szolgálnak a különböző regularizációs technikák, például L2 regularizáció, dropout, adataugmentáció vagy az adversarial training. A domain adaptation módszerek segítik a modelleket abban, hogy új környezetben is megbízhatóan működjenek.

2.2. Mélytanuló architektúrák katonai alkalmazásokhoz

A mélytanuló architektúrák jelentik a modern katonai MI-rendszerek alapját. A konvolúciós neurális hálózatok (CNN) elsősorban a vizuális információk feldolgozásában játszanak szerepet. Műholdképek, UAV-felvételek vagy földi kamerák adatai alapján képesek felismerni járműveket, fegyvereket, infrastruktúrákat vagy mozgó célpontokat. A hadszíntér azonban gyakran olyan körülményeket teremt, amelyek megnehezítik a vizuális felismerést. A rossz fényviszonyok, az álcázás, a terep változékonysága vagy az ellenséges fél által alkalmazott megtévesztő technikák mind csökkenthetik a CNN-ek teljesítményét. További probléma a betanító adatokban lévő zaj megjelenése. Alapvetően eltérő megközelítést igényel egy katonai célra szánt gépi látáson alapuló MI rendszer betanítása, mint például egy hasonló, de polgári használatra szánt rendszeré. Az egyik legalapvetőbb különbség már a betanító adatokban jelentkezik. A katonai gépi látás tanítása alapvetően eltér a polgáritól: a lerombolt infrastruktúra, az ellenséges, vagyis nem együttműködő

³ Berner et al. 2020

⁴ Nitzl et al. 2024

⁵ Das et al. 2024

⁶ Bowen–Jie, 2021, 22

⁷ Rahimi–Alahi, 2024

⁸ Pande–Wu [é. n.]

környezet és a baráti csapatok életvédelmi prioritása⁹ olyan extrém zajt és domain shift-et eredményez, amelyre a polgári modellek nem alkalmasak.

A transformer architektúrák a természetes nyelvfeldolgozás területén hoztak áttörést, de mára a katonai döntéstámogatásban is meghatározóvá váltak. Képesek nagy mennyiségű szöveges adatot feldolgozni, például hírszerzési jelentéseket, kommunikációs mintákat vagy nyílt forrású információkat. A többmodalitású modellek lehetővé teszik a képi és szöveges információk együttes értelmezését, ami jelentősen növeli a szituációs tudatosságot. Az ilyen integrátor szerepet betöltő egyik legismertebb rendszer a Palantir Ontology System, amely többféle forrásból dolgozó automatizációs és döntéstámogató platform.¹⁰

A gráf alapú hálózati struktúrák, vagyis a Graph Neural Networks-ök (GNN) különösen alkalmasak olyan katonai feladatokra, ahol a rendszer természeténél fogva gráfstruktúrájú. Ilyen például a logisztikai hálózatok optimalizálása, a kommunikációs hálózatok sérülékenységeinek feltárása vagy a drónrajok koordinációja. A GNN-ek képesek a hálózat szerkezetéből következtetni a rendszer viselkedésére, ami stratégiai előnyt jelenthet.¹¹

A megerősítéssel tanuló autonóm rendszerek egyik legfontosabb technikája. A drónok, robotok vagy elfogórendszerek képesek a környezet visszajelzései alapján megtanulni, hogyan hajtsanak végre komplex feladatokat. A jutalmazási struktúra meghatározza, hogy a rendszer milyen viselkedést tart kívánatosnak. A katonai környezet azonban nem determinisztikus, az ellenséges fél adaptív viselkedése pedig tovább nehezíti a tanulást. A multi-agent reinforcement learning lehetővé teszi, hogy több autonóm egység együttműködve, decentralizált módon hozzon döntéseket. Ez a drónrajok működésének alapja, ugyanakkor instabilitást és nem kívánt emergens viselkedést is eredményezhet. Másodlagos problémaként a különböző autonóm egységek szerepeinek kiosztása merül fel. Teljesen más módon szükséges betanítani egy olyan rendszert, ahol minden egység birtokában van a raj teljes tudásának és egy olyan rendszert, ahol vannak egyedi szerepkörök és ezek együttműködésére van szükség. Előbbi előnye, hogy bármely egység kiesése esetén annak szerepét a többi drón feloszthatja egymás között. Hátránya az egyes részfeladatok generalizálásának nehézsége. Egy specializált feladatra történő betanításhoz képest hasonló eredményt elérni nagyobb hálózattal lehetséges, amely az egyes eszközök esetleges limitált számítási kapacitás miatt komoly kihívást jelent.¹²

2.3. Adat- és szenzorarchitektúrák

A hadszíntér környezete általában nem teszi lehetővé, hogy minden adatot központi szerverekre továbbítsanak. A kommunikációs csatornák sérülékenyek, a sávszélesség korlátozott, az ellenséges fél aktív zavaró tevékenységet végez. Ezért a modern MI-rendszerek hierarchikus, többszintű feldolgozási modellt alkalmaznak, amely három fő rétegből áll: központi, regionális és lokális feldolgozásból.¹³

A központi adatfeldolgozás felel a stratégiai elemzésért, de nagy késleltetéssel tud dolgozni. A regionális feldolgozás a taktikai adatfúzió terepe, de megjelenik a redundancia. A lokális, edge adatfeldolgozás platformszintű, valós idejű döntéshozatalt biztosít. Ezen felül biztosítja a működést, valamint a feladatvégrehajtást például kommunikációs zavarok esetén is. Teljesen autonóm rendszerek esetén az ölési lánc technikai végrehajtása fedélzeti szinten történik.¹⁴

⁹ Padányi, 2017

¹⁰ Palantir: Palantir Ontology Overview 2025

¹¹ Yang et al. 2023

¹² Liao et al. 2025

¹³ Clark et al. 2020

¹⁴ 3000.09 DoD irányelv

A háromszintű feldolgozási modell stratégiai jelentőséggel bír. Lehetővé teszi a decentralizált döntéshozatalt, növeli a rendszer túlélőképességét és csökkenti a kommunikációs zavarok hatását. A döntéshozatal helye azonban jogi és etikai kérdéseket is felvet, mivel meghatározza, hogy milyen mértékben szükséges emberi felügyeletet biztosítani.

A katonai adatok minősége gyakran problémás. A szenzorok hibázhatnak, az adatok hiányosak lehetnek, a csapatok tévedhetnek a jelentések során, és az ellenséges fél aktív manipulációs kísérletei is torzíthatják a bemeneteket. Az adversarial machine learning (AML) azokat a tevékenységeket jelenti, amikor az ellenséges fél aktívan manipulálja a bemeneteket, például apró perturbációkkal, amelyek az ember számára észrevehetetlenek, de a modell számára félrevezetőek. A rendszerek azonban nemcsak a kifinomult perturbációkra, hanem a közvetlen jelzavarásra vagy a környezeti adatok szisztematikus hamisítására is érzékenyek. Utóbbira jó példa a GNSS spoofing, míg a jelzavarás esetén például a radarzavarásra gondolhatunk. A védekezés a kiberhabviseléshez hasonlóan alapvetően függ a támadás típusától, de az általánosan kijelenthető, hogy egy ellenállóbb rendszer eléréséhez robusztus tanulási módszerekre, redundáns szenzorokra és folyamatos ellenőrzésre van szükség. A bizonytalanság számszerűsítése lehetővé teszi, hogy a rendszer ne csak egyetlen kimenetet adjon, hanem annak megbízhatóságát is jelezze, ami a katonai döntéshozatalban különösen fontos.

3. A katonai döntéshozatal átalakulása: OODA és OODA(B)

A katonai döntéshozatal egyik legismertebb és legszélesebb körben alkalmazott modellje John Boyd ezredes OODA-ciklusa, ahogyan arra Rule¹⁵ rámutat. Ez a ciklus a megfigyelés, orientáció, döntés és cselekvés négyesére épül. A modell eredetileg az emberi döntéshozatal gyorsaságának és adaptivitásának leírására szolgált gyorsan változó és információhiányos környezetben. A modern hadviselésben azonban a döntéshozatal egyre nagyobb részét az MI rendszerek végzik, amelyek működése alapvetően eltér az emberi gondolkodástól. Ez indokoltá teszi a klasszikus OODA-modell kiterjesztését. A PDCA ciklussal szemben, amely fázisai a tervezés, cselekvés, ellenőrzés és beavatkozás, az OODA jobban illeszkedik az MI rendszerek folyamataihoz, mert nem szükséges idő előtti elköteleződés. A cselekvés a megelőző adatgyűjtésen alapul, amely fontos eleme a gépi tanulós rendszereknek. Rule munkájában¹⁶ idézi Boyd-ot, aki szerint a gépek nem vívnak háborút. Az emberek harcolnak és ők használják az eszközt. Ezen idézet fényében a korábbi kijelentések egészen érdekessé, már-már ironikussá válnak.

A modell paraméterei a betanítási fázis során a környezetből érkező adatok alapján folyamatosan frissülnek, a rendszer teljesítménye iteratív módon javul. Ez a folyamat matematikailag a gradiensalapú optimalizációval írható le, amely a gépi tanulás egyik alapmechanizmusa. Természetesen az ritkán fordul elő, hogy valós környezetben, az inferencia szakaszban visszatanítás is történik, de a begyűjtött adatok egy későbbi tanulási ciklusban felhasználásra kerülnek. A klasszikus OODA-ciklus azonban nem tartalmaz olyan elemet, amely a tanulási folyamatot vagy a visszacsatolt paraméterfrissítést önálló elemként reprezentálná. Kutatásom során ezért arra jutottam, szükségessé válik egy új komponens bevezetése, amely a modern MI-rendszerek működését pontosabban írja le.

A kiterjesztett OODA(B) modellben a „B” a backpropagation rövidítése, amely a visszacsatolt tanulást vagy visszatanítást jelenti. A backpropagation nem csupán egy technikai részlet, hanem a rendszer adaptivitásának és önjavító képességének kulcsa. A modell így öt elemből áll: megfigyelés, orientáció, döntés, cselekvés és visszacsatolt tanulás. A „B” komponens beillesztése lehetővé teszi, hogy a döntéshozatali ciklus ne csak végrehajtási folyamatként, hanem tanulási

¹⁵ Rule, 2013

¹⁶ Rule, 2013 i.m. 9.

folymatként is értelmezhető legyen. Ezzel a kiegészítéssel a ciklus nem csak lefedi a gépi tanulós döntéshozatalt, hanem reprezentálja egy MI rendszer teljes működését.

Az OODA(B) modellben a visszacsatolt tanulás tehát nem egyszerűen a ciklus végén jelenik meg, hanem a teljes folyamatot áthatja. Egy teljesen adaptív megoldás esetén például a rendszer minden döntés után értékeli a kimenetet és a hibák alapján módosítja a belső paramétereit. Ez a mechanizmus teszi lehetővé, hogy az MI alkalmazkodjon a környezet változásaihoz, felismerje a mintázatok, javítsa a teljesítményét. Csak a gépeknek adott autonómia foka és a felelősség kezelésének szabályozása szab határt a lehetséges implementációknak. Ezen logika alapján építhetők anomália-detekciós algoritmuson alapuló döntéshozatali rendszerek, vagy olyan megerősítéses tanulós folyamatok, amelyek a teljes súlyozás megváltoztatására is képesek. Itt fontos hangsúlyozni, hogy ez csupán egy technológiai lehetőség, a katonai alkalmazás rengeteg jogi kérdést felvet.

A modell mikro-szinten leírja egy-egy MI-rendszer működését egy konkrét feladat során. Egy autonóm drón például megfigyeli a környezetét, értelmezi az adatokat, döntést hoz a mozgásáról, végrehajtja azt, majd a visszacsatolt tanulás révén javítja a navigációs modelljét. A folyamat minden ciklusa során finomodik a rendszer viselkedése, ami hosszú távon nagyobb pontosságot és megbízhatóságot eredményez. Fontos figyelni a határok megalkotására, hiszen a súlyok folyamatos változtatásával bár az eszközök egyre jobban adaptálódnak a környezethez, a generalizációs képességük könnyen csökkenhet. Ezen túlmenően a jutalmazási struktúra hibái tehát olyan viselkedést eredményezhetnek, amely rövid távon előnyös, de hosszú távon veszélyes. A modell túlzott önbizalma vagy a bizonytalanság alulbecslése eszkalációs kockázatot hordozhat. Ezért a mikro-szintű működés folyamatos szabályozottságot, felügyeletet és ellenőrzést igényel.

Makro-szinten az OODA(B) modell a teljes hadműveleti rendszer működését írja le. A drónrajok, autonóm járművek, kiber-MI rendszerek és ISR-láncok mind olyan egységek, amelyek folyamatosan tanulnak és alkalmazkodnak. A visszacsatolt tanulás a modell paramétereinek folyamatos frissítését biztosítja, ami javítja a rendszer alkalmazkodóképességét és közvetve támogatja a komplex mintázatok felismerését. A makro-szintű adaptivitás azonban kockázatokat is hordoz, mivel a rendszer viselkedése idővel megváltozhat, emergens mintázatok jelenhetnek meg, amelyek nem mindig kívánatosak. A decentralizált döntéshozatal eredményeként az egyes egységek viselkedése összeadódik és az apróbb hibák is eszkalációs spirált indíthatnak el.

Az OODA(B)-ciklus a két szintnek köszönhetően egyszerre két idősíkon is mozog, egy taktikai és egy tanulási idősíkon. Előbbin a gyors döntéshozatal és reakcióképesség van a középpontban, míg utóbbin a folyamat közben megszerzett adatok generalizált becsatornázása a modellfrissítésbe.

Az OODA(B) modell egyik legfontosabb előnye, hogy képes integrálni az MI tanulási folyamatait a katonai döntéshozatalba. A klasszikus OODA-ciklus és a kiterjesztett OODA(B) modell közötti különbség nem pusztán egy új elem hozzáadása, hanem szemléletváltás is. Előbbi modell az emberi döntéshozatal gyorsaságát és rugalmasságát hangsúlyozza, míg a kiterjesztett változat a gépi tanulás adaptivitását és iteratív fejlődését is figyelembe veszi. Ez a kettős szemlélet lehetővé teszi, hogy a katonai rendszerek egyszerre legyenek gyorsak és tanulóképesek, ami a modern hadviselésben jelentős előnyt jelent. A jövő háborújában ugyanis a döntéshozatal egyre nagyobb részét MI rendszerek végzik majd, amelyek működése nem írható le pusztán a megfigyelés, orientáció, döntés és cselekvés négyesével. Boyd idézetét parafrázálva, a jövőben nem emberek vívnak háborút. Gépek harcolnak és használják a mesterséges eszüket.

4. Etikai és stratégiai kérdések

Az MI katonai célú alkalmazása nem csupán technológiai vagy jogi kérdés, hanem mély etikai és stratégiai dilemmákat is felvet. A modern hadviselésben az MI olyan képességeket hoz létre,

amelyek alapjaiban változtatják meg a konfliktusok természetét, a döntéshozatal szerkezetét és a felelősség fogalmát. A technológiai fejlődés gyorsasága miatt a normák, szabályok és etikai keretek gyakran csak késéssel követik a változásokat, ami bizonytalanságot és kockázatot eredményez. A kérdések nagyon széles skálán mozognak, amelyek részletes megvitatása önálló kutatási munkát igényelne, így ebben a fejezetben kizárólag összefoglaló, példálózó jelleggel kerülnek bemutatásra az egyes kérdéskörök.

4.1. A belépési küszöb csökkenése

Az MI egyik legfontosabb stratégiai hatása, hogy csökkenti a fegyveres konfliktusok belépési küszöbét. A hagyományos hadviselésben a katonai műveletek megkezdése jelentős emberi, politikai és gazdasági kockázattal járt. Az autonóm rendszerek azonban lehetővé teszik, hogy egy állam úgy hajtson végre támadást vagy felderítést, hogy közben nem kockáztat emberi életet. Ez a változás csökkenti a politikai költségeket és olyan helyzeteket teremthet, ahol a döntéshozók könnyebben vállalnak katonai akciókat.

A belépési küszöb csökkenése különösen veszélyes olyan helyzetekben, ahol a konfliktusok gyorsan eszkalálódhatnak. Az autonóm rendszerek alkalmazása olyan döntéseket is lehetővé tesz, amelyek korábban túl kockázatosnak számítottak volna. A fegyveres konfliktusok valószínűségének növekedése olyan helyzeteket teremthet, ahol a katonai műveletek megkezdése nem jár kellő politikai mérlegeléssel. A hadviselés így könnyen eljuthat a gazdasági döntések szintjére.

4.2. Túlzott bizalom és automatizációs torzítás

Az MI rendszerek alkalmazása gyakran túlzott bizalmat eredményez a technológia iránt. A döntéshozók hajlamosak lehetnek feltételezni, hogy a rendszer objektív, pontos és megbízható, miközben a valóságban a modellek működése számos bizonytalanságot és hibalehetőséget tartalmaz.¹⁷ Az automatizációs torzítás azt jelenti, hogy az emberek hajlamosak elfogadni a rendszer döntéseit akkor is, ha azok hibásak vagy hiányos adatokon alapulnak. Katonai környezetben egy téves célazonosítás vagy hibás fenyegetésvértékelés súlyos következményekkel járhat, és eszkalációs spirált indíthat el. A túlzott bizalom csökkenti az emberi felügyelet hatékonyságát, ahol a döntéshozók nem kérdőjelezik meg a rendszer működését. Ezért nagyon fontos a kritikus gondolkodás és a rendszeres ellenőrzés.

4.3. A háború köde és a dehumanizáció

A háború köde a katonai műveletek egyik alapvető jellemzője, amely a bizonytalanságból, a hiányos információkból és a gyorsan változó környezetből fakad. Az MI rendszerek képesek csökkenteni a háború kódét azáltal, hogy nagy mennyiségű adatot dolgoznak fel és pontosabb képet adnak a helyzetről. Ugyanakkor a technológia új formában teremti meg a bizonytalanságot. A modellek működése gyakran átláthatatlan, a döntések okai nem mindig érthetőek, amelyek súlyos hibákhoz vezethetnek.¹⁸

A dehumanizáció kérdése szintén fontos etikai dilemma. Az autonóm rendszerek alkalmazása csökkenti az emberi jelenlétet a harctéren, ami távolságot teremt a döntéshozók és a következmények között, ami a háború kódének növekedését jelenti. A katonai műveletek így könnyebben válhatnak technikai problémává, ahol a döntések nem emberi élethez, hanem algoritmusokhoz kapcsolódnak. Ez a folyamat csökkentheti az empátiát és a felelősségérzetet, ami veszélyes irányba terelheti a katonai döntéshozatalt.

¹⁷ Horowitz–Kahn 2023

¹⁸ Drinkall, 2025

4.4. Hiperháború és a műveleti tempó gyorsulása

A hiperháború fogalma azt írja le, amikor az MI rendszerek olyan mértékben gyorsítják fel a katonai műveleteket, hogy az emberi döntéshozók már nem képesek követni az eseményeket. Ilyen események már megtörténtek olyan területeken, ahol az algoritmikus döntéshozatalnak nagyobb múltja van, például a tőzsdén. A 2010-es Flash Crash eseménykor az algoritmusok körülbelül 36 perc leforgása alatt hatalmas rendszerszintű veszteséget okoztak.¹⁹ A drónrajok, autonóm fegyverrendszerek és kiber-MI rendszerek olyan sebességgel képesek reagálni a környezet változásaira, amely meghaladja az emberi reakcióidőt. Ez a folyamat a döntéshozatalt részben vagy teljesen a gépek kezébe helyezheti, amennyiben az emberi ellenőrzés nem kellőképpen biztosított.

A hiperháború egyik legnagyobb veszélye az eszkaláció. Ha a rendszerek automatikusan reagálnak a fenyegetésekre, illetve a döntéshozatal decentralizált, akkor olyan helyzetek alakulhatnak ki, ahol a konfliktusok emberi szándék nélkül is súlyosbodnak. A műveleti tempó gyorsulása csökkenti a diplomáciai és politikai beavatkozás lehetőségét, így a konfliktusok kontrollálhatatlanná válhatnak.

5. A katonai mesterséges intelligencia és a hadijog

A háborúval kapcsolatos normák hagyományosan négy nagy kategóriába rendeződnek: a jus contra bellum az erőszak általános tilalmát és az önvédelem feltételeit szabályozza, a jus ad bellum a fegyveres erő alkalmazásának jogszerűségét vizsgálja, a jus in bello a fegyveres konfliktusok során alkalmazandó humanitárius normákat határozza meg, míg a jus post bellum a háború lezárását és a felelősség kérdését rendezi. Az MI megjelenése mindegyik területet új kihívások elé állítja, mivel olyan döntéshozatali és végrehajtási mechanizmusokat hoz létre, amelyekre a jelenlegi jogi keretek nem készültek fel. Általánosan elmondható, hogy a fejezetben tárgyalt témakörök és a súlyos jogsértések elkerülésének szempontjából kulcsfontosságú, hogy az autonóm rendszerek működése átlátható legyen és a döntéshozatalban megmaradjon az emberi felügyelet, a megfelelően kimunkált felelősségi alakzatokkal együtt.

5.1. Jus contra bellum

A jus contra bellum a nemzetközi jog egyik alapköve, amely kimondja az erőszak alkalmazásának általános tilalmát és csak két kivételt ismer: az ENSZ Biztonsági Tanácsának felhatalmazását és az önvédelem jogát.²⁰ Az erőszak jogszerűségének megítélése nehézkes az MI szemüvegén keresztül.

A prediktív elemzésre épülő rendszerek képesek előre jelezni az ellenséges mozgásokat vagy fenyegetéseket, de ezek a predikciók gyakran bizonytalanok és hibalehetőségeket tartalmaznak. Ha egy állam egy MI által jelzett fenyegetés alapján indít katonai műveletet, felmerül a kérdés, hogy ez megfelel-e az önvédelem jogának. A jelenlegi jogi értelmezés szerint az önvédelem *conditio sine qua non*-ja az agresszió vagy fegyveres támadás megtörténte.²¹ A preventív önvédelem kiterjesztő értelmezése szerint az önvédelem akkor is jogszerű, ha a támadás közvetlenül fenyeget. A nemzetközi joggyakorlat feladata annak kimunkálása a jövőben, hogy az MI által jelzett valószínűségi indikáció minősülhet-e közvetlen fenyegetésnek. A jus contra bellum szempontjából ezért kulcsfontosságú, hogy az MI rendszerek működése átlátható legyen, a döntéshozatalban pedig megmaradjon az emberi kontroll.

¹⁹ Kirilenko et al. 2017

²⁰ Bruhács 2014, 71.

²¹ Uo. 76.

5.2. Jus ad bellum

A jus ad bellum azt vizsgálja, hogy egy állam mikor alkalmazhat fegyveres erőt jogszerűen. Ebben a kontextusban kétféle kihívás jelenik meg. Egyrészt a prediktív elemzésre épülő rendszerek olyan fenyegetéseket is azonosíthatnak, amelyek valójában nem léteznek. Másrészt az autonóm rendszerek működése olyan gyors, hogy a döntéshozók nem mindig képesek megfelelően mérlegelni a helyzetet.

A megelőző csapás kérdése különösen problémás, mert a nemzetközi jog jelenlegi értelmezése szerint ez csak akkor jogszerű, ha a támadás közvetlenül fenyeget. Az MI által jelzett valószínűségi indikáció azonban nem minősül közvetlen fenyegetésnek, ezért az ilyen alapon indított katonai művelet jogszerűsége megkérdőjelezhető. Az MI tehát nem szolgálhat önmagában jogalkapításként a fegyveres erő alkalmazására.²²

A jus ad bellum szempontjából az is kérdés, hogy az autonóm rendszerek által végrehajtott cselekmények kinek tulajdoníthatók. Ha egy autonóm rendszer tévesen azonosít egy fenyegetést, majd támadást indít, felmerül a kérdés, hogy ez az állam cselekményének minősül-e. A jelenlegi jogi keretek szerint az állam felel a területén vagy ellenőrzése alatt működő rendszerekért, függetlenül attól, hogy azok autonóm módon működnek-e.²³

5.3. Jus in bello

A jus in bello a fegyveres konfliktusok során alkalmazandó humanitárius normákat határozza meg. A legfontosabb elvek a megkülönböztetés, az arányosság és a szükségesség. Az autonóm fegyverrendszerek alkalmazása azonban komoly kérdéseket vet fel ezekkel az elvekkel kapcsolatban.

A megkülönböztetés elve azt követeli meg, hogy a katonai műveletek során különbséget kell tenni a katonai és civil célpontok között. A gépi látás rendszerei azonban nem mindig képesek pontosan azonosítani a célpontokat, különösen olyan környezetben, ahol az ellenséges fél aktív álcázást vagy megtévesztést alkalmaz. A gyakorlatban a domain shift, az adversarial támadások és a szenzorhibák mind csökkentik a rendszer pontosságát, ami sértheti a megkülönböztetés elvét.²⁴

Az arányosság elve azt követeli meg, hogy a katonai előny és a civil károk között megfelelő egyensúly legyen. Az autonóm rendszerek esetén ez a mérlegelés különösen nehéz, hiszen itt a rendszer célja a hatékony katonai akció, amit az elvnek felül kellene írnia. A kérdés tehát, hogy egy algoritmus képes lesz-e értékelni a járulékos károk morális súlyát (szubjektív mérlegelés), vagy csak matematikai optimalizációt végez.

A szükségesség elve alapján a támadás csak akkor jogszerű, ha nincs más alternatíva. Az autonóm rendszerek működése azonban gyakran automatizált és nem mindig veszi figyelembe a nem halálos vagy kevésbé káros alternatívákat.²⁵

5.4. Jus post bellum

A jus post bellum a háború lezárását és a felelősség kérdését szabályozza. Az autonóm rendszerek alkalmazásánál, különösen hibás döntések meghozatala esetén a felelősség megállapítása

²² Grimal–Pollard 2021

²³ Amoroso 2017, 29.

²⁴ Uo. 13.

²⁵ Uo. 14.

különösen nehéz. A fejlesztők, a gyártók, a parancsnokok és az operátorok mind szerepet játszanak a rendszer működésében, de a döntés végső soron a rendszer autonóm működésének eredménye. Természetesen nem autonóm működés esetén a felelősség megállapítása egyértelmű, a hatályos nemzetközi jogi dogmatika és felelősségi rezsimek szerint alakul, így követi a bevett nemzetközi joggyakorlatot.

A felelősség megállapítása mellett a helyreállítás kérdése is fontos. Az MI által okozott károk gyakran nehezen rekonstruálhatók, mivel a rendszerek működése nem mindig átlátható. A visszacsatolt tanulás miatt a rendszer viselkedése idővel megváltozhat, ami megnehezíti a döntések okainak feltárását. Ezért a *jus post bellum* szempontjából kulcsfontosságú, hogy az autonóm rendszerek működése auditálható legyen, és rendelkezzenek olyan mechanizmusokkal, amelyek lehetővé teszik a döntések visszakövetését.

6. Összegzés, a jövő irányai

Az MI katonai alkalmazása nem csupán új eszközöket hoz létre, hanem átalakítja a hadviselés egész szerkezetét. A technológiai fejlődés gyorsasága, a döntéshozatal automatizálása és a rendszerek adaptív működése olyan új dinamikákat hoz létre, amelyek a konfliktusok természetét is megváltoztatják. A jövő hadviselésének megértéséhez olyan fogalmi keretre van szükség, amely képes egyszerre kezelni a technológiai, jogi, stratégiai és etikai dimenziókat. A jövő hadviselésének stratégiai keretei öt olyan alapfogalom mentén határozhatók meg, amelyek az MI által formált hadviselés legfontosabb területeit és irányait jelölik ki: predikció, prevenció, eszkaláció, katalizáció és redempció.

6.1. Predikció: a jövő előrejelzése mint katonai képesség

A modern MI-rendszerek képesek hatalmas mennyiségű adatot feldolgozni és olyan mintázatokat felismerni, amelyek az emberi elemzők számára rejtve maradnának. A predikció a hadviselés minden szintjén megjelenik: a csapatmozgások előrejelzésétől a logisztikai igények becslésén át a kiberfenyegetések azonosításáig. A prediktív képességek azonban nem mentesek a bizonytalanságtól. A modellek működése valószínűségi alapon történik, a hibák pedig súlyos következményekkel járhatnak. Egy téves előrejelzés eszkalációt indíthat el, különösen akkor, ha a döntéshozók túlzott bizalmat helyeznek a rendszer működésébe. A predikció ezért nem csupán technológiai, hanem stratégiai és jogi kérdés is, amely egyszerre hordoz lehetőségeket és komoly kockázatokat. A jövő hadviselésében azok az államok lesznek előnyben, amelyek képesek a prediktív képességeket felelősen, átláthatóan és jogszerűen alkalmazni.

6.2. Prevenció: a konfliktusok megelőzése és a korai beavatkozás dilemmái

A predikció logikus következménye a prevenció. Ha a rendszer képes előre jelezni a fenyegetéseket, akkor lehetőség nyílik azok megelőzésére vagy korai kezelésére. A prevenció a modern hadviselés egyik legfontosabb célja, mivel a konfliktusok elkerülése vagy korai kezelése jelentős politikai, gazdasági és emberi előnyökkel jár. Az MI azonban elmoshatja a prevenció és a megelőző csapás közötti határt, ahol a megelőző csapás logikusabb és hatékonyabb válaszlépésnek tűnik, mint a több energiával járó megelőzés. A prevenció etikai kockázata, hogy egy téves gépi riasztás emberi szándék nélküli, automatikus konfliktushoz vagy eszkalációhoz vezethet.

6.3. Eszkaláció: a gyorsuló hadviselés kockázatai

Az autonóm rendszerek működése olyan gyors, hogy az emberi döntéshozók gyakran nem képesek követni az eseményeket. A gyorsuló műveleti tempó csökkenti a diplomáciai és politikai beavatkozás lehetőségét, és olyan helyzeteket teremthet, ahol a konfliktusok kontrollálhatatlanná

válnak. Második értelmezés szerint amennyiben a konfliktushoz nem szükséges emberi életeket kockáztatni, mert rendelkezésre áll nagy mennyiségű gépi erőforrás, akkor sokkal könnyebben lehet a fegyveres konfliktus a végleges döntés. A rendszeres ellenőrzés, a kritikus gondolkodás és a jogi keretek betartása nélkül az eszkalációs kockázat jelentősen nő.

6.4. Katalizáció: az MI mint erősorzó és rendszerszintű átalakító tényező

Az MI nem csupán végrehajtja a feladatokat, hanem fel is erősíti a katonai képességeket, amelyek átalakítják a hadviselés egész szerkezetét. A gépi látás rendszerei növelik a felderítés hatékonyságát, a prediktív analitika javítja a döntéshozatalt, a megerősítéses tanulás pedig lehetővé teszi az autonóm rendszerek adaptív működését. A katalizáció azonban nem csupán technológiai előnyt jelent, hanem stratégiai kockázatot is. Az MI alkalmazása csökkenti a belépési küszöböt, növeli a műveleti tempót és decentralizálja a döntéshozatalt. Ezek a folyamatok olyan helyzeteket teremthetnek, ahol a konfliktusok gyorsabban és kiszámíthatatlanabban alakulnak. A katalizáció ezért egyszerre jelent lehetőséget és kockázatot, amelyet felelősen kell kezelni.

6.5. Redempció: felelősség, helyreállítás és a háború utáni igazságosság kérdése

A redempció a háború utáni igazságosság fogalmához kapcsolódik. A fogalom itt a technológiai és morális rehabilitációt jelenti, tehát nem a szakrális értelemben vett megváltást vagy bűnbocsánatot jelenti. Az MI által okozott károk helyreállítása, a felelősség megállapítása és a bizalom helyreállítása olyan feladatok, amelyek különösen összetettek. Felmerül a kérdés, hogy az a rendszer, amelyik a károkat okozta, képes lehet-e azokat enyhíteni. Extrém esetben például megbíznának-e az emberek abban a járőröző, rendőri feladatokat ellátó drónban, amelyik hetekkel korábban még embereket ölt? E kérdésben nem csak maga a szoftver számít, hanem az eszköz megjelenési formája, vagy éppen a gépek általános megítélése is. A bizalomvesztés ugyanis nem szükségszerűen csak az adott verziójú MI szoftvert fogja érinteni, hanem az egész rendszert, benne mindenféle gépi tanuláson alapuló eszközzel.

Irodalomjegyzék

Amoroso, Daniele (2017): *Jus in bello and jus ad bellum arguments against autonomy in weapons systems: A re-appraisal*. QIL 43 Questions of International Law, Zoom-in 5

Berner, Julius et al. (2020): *Analysis of the Generalization Error: Empirical Risk Minimization over Deep Artificial Neural Networks Overcomes the Curse of Dimensionality in the Numerical Approximation of Black-Scholes Partial Differential Equations*. Arxiv, DOI: 10.48550/arXiv.1809.03062

Bowen, Yu – Jie, Zhang (2021): *Military Target Detection Method Based on Improved YOLOv3 Network*. International Journal of Engineering and Applied Science, 8. évf., 2021/9, 17–24.

Bruhács János (2014): *Jus contra bellum – Glosszák az erőszak nemzetközi jog tilalmához*. In: Blutman László (szerk.): Ünnepi kötet dr. Bodnár László egyetemi tanár 70. születésnapjára. Szegedi Tudományegyetem Állam- és Jogtudományi Kar, Szeged, 69-84.

Clark, Bryan et al. (2020): *Mosaic warfare exploiting artificial intelligence and autonomous system to implement decision-centric operations*. Center for Strategic and Budgetary Assessment, https://csbaonline.org/uploads/documents/Mosaic_Warfare_Web.pdf (Letöltés időpontja: 2026. 01. 29.)

Das, Gautom et al. (2026): *Towards Understanding Best Practices for Quantization of Vision-Language Models*, Arxiv, DOI: doi.org/10.48550/arXiv.2601.15287

DoD (2023): 3000.09. 2023.01.05-i irányelve a fegyverrendszerek autonómiájáról (*Autonomy In Weapon Systems*)

Drinkall, Toby (2025): *Red Lines and Grey Zones in the Fog of War: Benchmarking Legal Risk, Moral Harm, and Regional Bias in Large Language Model Military Decision-Making*. Arxiv, DOI: 10.48550/arXiv.2510.03514

Grimal, Francis – Pollard, Michael J (2021): *Embodied artificial intelligence and jus ad bellum necessity: Influence and imminence in the digital age*. Geo. J. Int'l L., Vol 53, 2021, 209-275.

Horowitz, Michael C. – Kahn, Lauren (2023): *Bending the Automation Bias Curve: A Study of Human and AI-based Decision Making in National Security Contexts*. Arxiv, DOI: 10.48550/arXiv.2306.16507

Kirilenko, Andrei et al. (2017): *The Flash Crash: The Impact of High Frequency Trading on an Electronic Market*. CFTC, DOI: 10.2139/ssrn.1686004

Liao, Xiaomin et al. (2025): *Heterogeneous Multi-Agent Deep Reinforcement Learning for Cluster-Based Spectrum Sharing in UAV Swarms*. MDPI, DOI: 10.3390/drones9050377

Nitzl, Christian et al. (2024): *The Use of Artificial Intelligence in Military Intelligence: An Experimental Investigation of Added Value in the Analysis Process*. Arxiv, DOI: 10.48550/arXiv.2412.03610

Padányi József (2017): *Egy kínai hadtudós gondolatai (Szun-Ce: A hadviselés törvényei)*. In: Gőcze István (szerk.): Állam és Katona, Dialóg Campus Kiadó, Budapest

Palantir [@palantirtech] (2025): *Palantir Ontology Overview*. Youtube, 2025. 11. 17. <https://www.youtube.com/watch?v=YDAxITCNcko> (Letöltés időpontja: 2026. 01. 30.)

Pande, Swapnil – Wu, Xindi (é.n.): *Marrying Motion Forecasting and Offline Model-Based Reinforcement Learning for Self-Driving Cars*. [k. n.], <https://xindiwu.github.io/data/motion.pdf>, (Letöltés időpontja: 2026. 01. 30.)

Rahimi, Ahmad – Alahi, Alexandre (2024): *A Multi-Loss Strategy for Vehicle Trajectory Prediction: Combining Off-Road, Diversity, and Directional Consistency Losses*. Arxiv, DOI: 10.48550/arXiv.2411.19747

Rule, Jeffrey N (2013): *A Symbiotic Relationship: The OODA Loop, Intuition, and Strategic Thought*. U.S. Army War College, Pennsylvania 17013

Vassilev, Apostol et al. (2024): *Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations*. NIST AI 100-2 E2023, Nist, DOI: 10.6028/NIST.AI.100-2e2023

Yang, Chenxiao et al. (2023): *Graph Neural Networks are Inherently Good Generalizers: Insights by Bridging GNNs and MLPs*. Arxiv, DOI: 10.48550/arXiv.2212.09034

Mesterséges intelligencia a levegőben

Artificial Intelligence in the Air

Katona Róbert¹ 

Absztrakt

A tanulmány célja annak bemutatása, hogy a dróntechnológia és a mesterséges intelligencia (MI) integrációja miként alakítja a különböző iparágak működését, valamint milyen technológiai, gazdasági és társadalmi változásokat eredményez. Különös figyelmet kap a modern mezőgazdasági termelés és a vizuális tartalomgyártás területe, ahol az innovatív, diszruptív technológiák új munkafolyamatokat és költséghatékonyabb megoldásokat teremtenek.

A tanulmány kiemelten vizsgálja az MI-alapú funkciókat, többek között az intelligens leszállási képességeket, valamint a filmiparban alkalmazott automatikus drónalapú kameravezérlést. Az eredmények szerint az MI-vel támogatott drónrendszerek jelentősen növelik a műveletek hatékonyságát és biztonságát, miközben csökkentik az emberi beavatkozás szükségességét és az üzemeltetési költségeket a hagyományos technológiákhoz képest.

A tanulmány rámutat arra is, hogy a technológia fejlődése hosszú távon új szakmák megjelenését eredményezi, ezáltal pedig a szakértői támogatás szerepe is felértékelődik.

Kulcsszavak: drón, mesterséges intelligencia, filmipar, precíziós mezőgazdaság, VTOL

Abstract

The aim of this study is to demonstrate how the integration of drone technology and artificial intelligence (AI) influences the operations of various industries, and how it generates technological, economic, and social changes. Special attention is given to modern agricultural production and content creation, where innovative and disruptive technologies create new workflows and more cost-efficient solutions.

The study focuses on AI-based functionalities, including intelligent landing capabilities, and automated drone-based camera control in the film industry. The results indicate that AI-supported drone systems significantly enhance operational efficiency and safety, while reducing the need for human intervention and operational costs compared to conventional technologies.

The study also highlights that technological advancements will lead to the emergence of new professions in the long term, thereby increasing the importance of professional support.

Keywords: drone, artificial intelligence, film industry, precision agriculture, VTOL

¹ IT Manager, katona.robert.info@gmail.com

1. Bevezetés

A dróntechnológia az elmúlt évtizedben jelentős fejlődésen ment keresztül, és napjainkra széles körben alkalmazott, meghatározó technológiai eszközzé vált. Manapság egyre több piaci szereplő ismeri fel a dróntechnológiában rejlő lehetőségeket, különösen a korszerű mezőgazdasági permetezési megoldások területén. A precíziós mezőgazdaság elterjedésével egyidejűleg folyamatosan növekszik az igény a különböző funkciókkal rendelkező, széles körben alkalmazható légitjárművek iránt. Ezek az eszközök nem kizárólag a mezőgazdasági feladatok hatékonyabb elvégzését támogatják, hanem számos más iparágban is megjelennek, mint például a filmiparban, ahol új technikai és vizuális megoldásokat tesznek lehetővé.

2050-re a globális népesség jelentős növekedése várható, melynek következtében megnő az élelmiszerigény. E kihívások kezelése érdekében a mezőgazdaság modernizációja elengedhetetlen.²

Az egyik jelentős technológiai előrelépést a dróntechnológia alkalmazása jelenti, különösen a permetező drónok alkalmazása amelyek a vegyszerek precízebb kijuttatása révén számos potenciális előnyt biztosítanak a hagyományos módszerekkel szemben.^{3 4}

A drónok kezdetben főként hobbi célokat szolgáltak és manuális irányítással működtek, napjainkra azonban egyre nagyobb szerepet kapnak a fejlett, mesterséges intelligenciával támogatott megoldások. Az MI integrációja lehetővé tette az autonóm és félautonóm működést, amely jelentősen növeli a drónok hatékonyságát, pontosságát és alkalmazhatóságát.

A hagyományos repülőeszközökhöz viszonyítva a pilóta nélküli légi járművek számos műszaki és gazdasági előnnyel rendelkeznek. Ezek közé sorolható a kisebb fizikai méret, a rugalmas felhasználhatóság, valamint az üzemeltetéshez kapcsolódó alacsonyabb beruházási, gyártási és fenntartási költségek.

A pilóta nélküli repülőeszközök többféle néven is ismertek a szakirodalomban, ezek közül az egyik leggyakrabban használt elnevezés az UAV, vagyis az Unmanned Aerial Vehicle. A tudományos és technológiai innovációk előrehaladásával párhuzamosan az UAV-rendszerek egyre összetettebb eszközökkel egészülnek ki, beleértve a különböző szenzortechnológiákat és a lézeres távérzékelési megoldásokat. Ennek eredményeként a polgári alkalmazások köre folyamatosan bővül, és az UAV-k egyre jelentősebb szerepet töltenek be számos iparágban.⁵

A tanulmány célja bemutatni, hogy a mesterséges intelligenciával támogatott dróntechnológiák milyen technológiai, gazdasági és társadalmi hatásokkal járnak. Az elemzés két kiemelt alkalmazási területre összpontosít: a mezőgazdasági precíziós drónmegoldásokra és a filmipar kreatív, költséghatékony technológiáira. Különös figyelmet kap a VTOL képesség, amely lehetővé teszi a vertikális fel- és leszállást, és amely mesterséges intelligenciával felokosított vezérlési funkciókkal rendelkezik. Ez a megoldás növeli a repülés stabilitását és biztonságát. Ezen túl áttekintésre kerülnek az autonóm rendszerekkel szembeni elvárások és korlátok, például a gyors feladatvégrehajtás és a magas hatékonyság, amelyeket a környezeti feltételek, a szenzorok és a vezérlési algoritmusok teljesítménye befolyásol. Végül a tanulmány rávilágít arra, hogy az új technológiák nemcsak a drónok alkalmazási területeit bővítik, hanem a meglévő munkakörökre is hatással vannak, új képességeket és felelősségeket követelve a szakemberektől. Mindez egyúttal a szakértői és tanácsadói támogatás szerepének felértékelődését is eredményezi. A dolgozat

² Vlaskality et al., 2025

³ Talaeizadeh et al., 2023

⁴ García-Munigua et al., 2024

⁵ Tsouros et al., 2019

átfogó képet nyújt arról, hogyan alakítja át a VTOL képesség és a mesterséges intelligencia a drónok biztonságosabb mindennapi alkalmazását.

2. MI fogalma

A mesterséges intelligencia napjaink egyik legdinamikusabban fejlődő technológiai területe, amely alapvetően átalakítja az ipari, szolgáltatói és kreatív szektorokat. A mezőgazdaságtól a filmiparig mindenhol megtalálható. Jól szemlélteti alkalmazásának sokszínűségét, hogy a mobiltelefonok virtuális asszisztenseiben és a különböző okoseszközökben is előfordul.

Jellemzően olyan számítógépes rendszerekről van szó, amelyek képesek bizonyos emberi kognitív folyamatok utánzására. Ezek a rendszerek képesek adatok elemzésére, mintázatok felismerésére, tanulásra és döntéshozatalra egyaránt.⁶

Általában elvárás velük szemben, hogy emberi beavatkozás nélkül reagáljanak a környezetből érkező információkra úgy, mintha egy emberi intelligenciával rendelkező lény hozna döntést.

Jelenleg nincs olyan mesterséges intelligencia definíció, amelyet a tudományos közösség általánosan elfogadna. Míg a számítástechnika gépi látásra vagy beszédfelismerésre koncentrálnak, más tudományterületek eltérő szempontokat és módszereket használnak, ezért jelenleg nincs általánosan elfogadott mesterséges intelligencia-definíció.⁷ Ennek következményeként a kutatók eltérő szempontok alapján értelmezik a mesterséges intelligencia fogalmát, és a meghatározás gyakran az adott alkalmazási területtől függ. A mezőgazdaságban kiemelt jelentőségű a precíziós permetezés, míg a filmiparban az AI hozzájárul az olcsóbb és minőségi felvételek elkészítéséhez.

3. Mesterséges intelligencián alapuló dróntechnológiákkal kapcsolatos elvárások

Bár az MI technológiák gyors fejlődése és elterjedése jelentős előnyöket biztosít az autonóm rendszerek és kreatív alkalmazások számára, ezzel párhuzamosan számos technikai, biztonsági és etikai kihívás is felmerül. Az autonóm MI-rendszerektől elvárnánk, hogy emberi beavatkozás nélkül képesek legyenek adaptívan reagálni a környezetből érkező információkra, önálló döntéseket hozni és feladatokat végrehajtani, ám a gyakorlatban ezt számos tényező korlátozza. Míg elméleti szinten az ideális autonóm rendszer teljesen függetlenül működik, a valós alkalmazásokban a legtöbbször szükség van bizonyos mértékű emberi felügyeletre.⁸

Ma már az autonóm működéstől nemcsak a műszaki megfelelőséget várjuk el, hanem azt is, hogy a rendszer valós környezetben megbízhatóan, biztonságosan és stabilan működjön.

Jelen tanulmány célja nem a drónok üzemeltetésére vonatkozó jogszabályok és engedélyek részletes ismertetése. Ezek tárgyalása az említés szintjén mégis fontos, mert a vonatkozó kodifikált előírások be nem tartása komoly jogi következményeket vonhat maga után. Az operátorok és felhasználók számára ezért az élet más területéhez hasonlóan elengedhetetlen, hogy tisztában legyenek az aktuális jogszabályokkal, és azokat maradéktalanul betartsák.

⁶ Hassani et al., 2020

⁷ Emmert-Streib et al., 2020

⁸ Zhang et al., 2020

4. Az autonóm rendszerekkel szembeni elvárások

Az autonóm drónrendszerek alkalmazásával kapcsolatban a piaci és felhasználói oldalon még nem alakult ki egységes és minden területre kiterjedő követelményrendszer. Ugyanakkor az igények és elvárások folyamatosan bővülnek.

Rejeb és társai szerint az elmúlt évtizedben jelentősen megnőtt a mezőgazdasági drónokról megjelent kutatások száma, ami az iparág fejlődésére utal.⁹ A permetezésre tervezett mezőgazdasági drónok esetében egyre több alkalmazási területen jelenik meg az a vélekedés, hogy a rendszerek önállóan hajtsák végre a kijuttatási műveleteket, optimalizálva a terület lefedettségét és a felhasznált növényvédő szer mennyiségét. Ez azt jelenti, hogy a permet minden szükséges növényhez eljusson, sehol ne maradjon ki terület, és a vegyszert a lehető legkisebb környezetterhelés mellett használják fel. Ugyanakkor az autonóm permetezésre használt drónrendszerek korlátai is egyértelműen érzékelhetők, különösen az új, korábban nem dokumentált növénybetegségek megjelenése esetén. Ilyenkor a rendszer felismerési pontossága csökkenhet, ami egyértelműen rámutat arra, hogy az emberi szakértelem és a folyamatos felügyelet továbbra is elengedhetetlen a működés során. Ennek megfelelően az autonóm döntéshozatal hatékonysága nagymértékben függ a rendszer betanításának minőségétől, az adatbázisok naprakészségétől, valamint az emberi beavatkozás lehetőségétől, amely lehetővé teszi a helyes reagálást olyan helyzetekben, amelyekre a mesterséges intelligencia még nem készült fel.

Az autonóm permetező drónoktól elvárt magas hatékonyság és gyors feladatvégrehajtás több tényező által is korlátozott. Az akkumulátorok kapacitása és az üzemidő például jelentősen meghatározza, hogy egy drón milyen hosszú ideig és milyen messzire tud repülni egy feltöltéssel. A jelenlegi technológiai korlátok miatt a mezőgazdasági drónokat gyakran a permetezési terület közelébe ajánlatos kitelepíteni, hogy a rendelkezésre álló energia a lehető legnagyobb mértékben ne a felesleges repülési útvonalakra hanem magára a permetezésre fordítódjon. Emellett a drón akkumulátorának cseréje vagy újra töltése, valamint a permetanyag pótlása is könnyebben és gyorsabban biztosítható, ha a drón a permetezési terület közelében van. Így a drón felügyeletét ellátó személy gyorsan beavatkozhat.

A mezőgazdasági drónok esetében a biztonságos fel- és leszállás alapvető elvárás. A hibák miatt megsérült eszközök javítása drága és időigényes, valamint a drónok nagy mérete miatt egy esetleges baleset során azok jelentős károkat okozhatnak a terményekben vagy az épített környezetben. A mesterséges intelligencia lehetővé teszi, hogy a drón valós időben feldolgozza a különböző szenzorok adatait, és gyorsan reagáljon a széllelőkések, dőlések és turbulenciák okozta eltérésekre. A mesterséges intelligencia hiányában a drón lassabban reagálna a környezet változásaira, így a kívánt pálya fenntartása gyorsan változó időjárási körülmények között jelentősen nehezebbé válna.¹⁰ A tanulmány rámutat arra, hogy a biztonságos fel- és leszállás az autonóm permetező drónok egyik legmeghatározóbb funkciója, mivel a technológia jelentős fejlődése és az iránta mutatkozó növekvő piaci kereslet egyaránt azt jelzi, hogy a felhasználók elsődlegesen megbízható és üzembiztos működést várnak el ezektől a rendszerektől. A rendszer alkalmazásának felelőssége a drón felügyeletét ellátó szakembert terheli, akinek feladata a környezeti feltételek folyamatos értékelése és a kockázatok mérlegelése. A drónpilóták felelőssége hasonlatos az önvezető autók sofőrjeihez, ahol a teljes felelősség a gépjármű sofőrjén van függetlenül az autó autonómiájának fokától. Amennyiben a meteorológiai viszonyok a biztonságos üzemeltetés határait meghaladják, a felügyelő időben beavatkozhat, és dönthet a műveletek felfüggesztéséről mindaddig, amíg a körülmények ismét alkalmassá nem válnak a

⁹ Rejeb et al., 2022

¹⁰ Caballero-Martín et al., 2024

biztonságos munkavégzésre. Ez is jól mutatja, hogy az autonóm technológia mellett az emberi jelenlét és felelősségvállalás továbbra is meghatározó tényező a kockázatok kezelésében.

A felhasználási környezet jelentősen befolyásolja, hogy az elvárások mennyire valósíthatók meg. Városi területeken vagy filmforgatások során számos jogi és biztonsági szempont merül fel, amelyek összetettebbek mint amik egy elhagyatott mezőgazdasági területen végzett művelet esetében megjelenik. Bár technikailag elvárható, hogy a drón képes legyen a feladatot végrehajtani, a jogszabályok és egyéb biztonsági előírások betartása a hatályos normák alapján a drónkezelő felelőssége. A filmiparban például gyakran elvárás, hogy a drón előre beprogramozott vagy betanított kameramozgásokkal stabilan és precízen rögzítsen jeleneteket. Ez azonban nem feltétlen csökkenti a filmes szakemberek szerepét. Ellenkezőleg, a technológia új vizuális lehetőségeket nyit meg, lehetővé téve komplexebb és dinamikusabb kameramozgásokat, amelyek közül sok korábban akár helikopterrel sem volt kivitelezhető. A drónok alkalmazása a filmiparban nem egyszerűen kiváltja az emberi munkát, hanem alapvetően átalakítja azt. A filmes szakembereknek, az operatőröknek, vágóknak és rendezőknek érdemes elsajátítaniuk a drónok üzemeltetését, hogy a technológia adta lehetőségeket kreatívan kihasználhassák. Az új kameramozgások dinamikus, korábban nem kivitelezhető perspektívákat tesznek lehetővé, de a vizuális hatásukat csak az emberi kreativitás révén lehet teljesen kiaknázni, például a vágás és kompozíció révén.

Az AI-alapú rendszerektől gyakran elvárjuk az önellenőrzési képességet is. A nagy mennyiségű szenzoradat feldolgozása lehetővé teszi bizonyos meghibásodások előrejelzését. Ez jelentősen növeli a repülésbiztonságot, azonban nem teszi feleslegessé a rendszeres karbantartást. A külső sérülések, mechanikai elhasználódások vagy egyéb, szenzorok által nem minden esetben észlelhető hibák továbbra is fizikai ellenőrzést igényelnek.

A felsorolt elvárások nem teljeskörűek de arra engednek következtetni, hogy az emberi tényező továbbra sem nélkülözhető. Az új munkakörök megjelenése illetve a meglévő munkakörök kibővítése új kihívásokat jelent a humánerőforrás-menedzsment számára. Ebben a folyamatban a tanácsadói szektor meghatározó szerepet tölthet be a szükséges kompetenciák fejlesztésében.

5. Az autonóm technológiák szervezeti hatásai és a tanácsadói szektor szerepe

Az autonóm technológiák megjelenése új szervezeti kihívásokat teremt, amelyek a munkafolyamatok és a döntési struktúrák átalakítását igénylik. Bár a technológia egyre nagyobb szerepet kap, az emberi tényező továbbra is meghatározó, hiszen a sikeres működéshez elengedhetetlen a dolgozók aktív bevonása és a meglévő tudás hatékony kihasználása.

A gyors technológiai fejlődés és az autonóm rendszerek alkalmazása új kihívások elé állítja a vállalkozásokat, különösen a szükséges szaktudás, tanácsadói támogatás, képzés és felelősségi körök kialakítása terén. A mesterséges intelligencia térnyerésével a munkaerőpiacon és a vizsgált iparágakban is változások figyelhetők meg. A kompetenciamentedzsment, az élethosszig tartó tanulás és a munkavállalói fejlesztés kulcsfontosságú szerepet kapnak a szervezetek sikeres működésében.¹¹

A HR és a tanácsadói szektor számára új kihívást jelentenek azok a változások, amelyek a digitális kompetenciák növekvő jelentőségét, valamint az alkalmazkodóképesség és rugalmasság készségeinek megerősödését igénylik. Különösen igaz ez a gyorsan változó munkaerőpiacon. A cégek olyan munkavállalókat és szakértői támogatást keresnek, akik képesek sikeresen navigálni az új technológiai kihívások között.

¹¹ Poór et al., 2025

6. Dróntechnológia előnyei a hagyományos módszerekkel szemben

A mezőgazdaságban az energia és környezeti hatékonyság kulcsfontosságú és a drónos technológia gyorsabb, környezetbarátabb alternatívát kínál a mezőgazdasági gépekkel szemben. Safaeinejad és társai kutatásukban kimutatták, hogy a drónos permetezés jelentősen energiatakarékosabb, mint a hagyományos traktoros módszer. A drón energiafelhasználása mindössze 146,84 MJ/ha, míg a traktoros permetezése 365,26 MJ/ha, ami azt jelenti, hogy a drón energiafelhasználása körülbelül 2,4-szer alacsonyabb. Az MJ, vagyis megajoule, az energia mérésére szolgáló egység, amely megmutatja, mennyi munka vagy energia szükséges egy adott feladathoz. A ha, vagyis hektár, a terület mértékegysége, amely 10 000 négyzetmétert jelent. A környezeti hatásokat vizsgálva a drónos permetezés során 14,485 kg CO₂/ha keletkezik, míg a hagyományos módszer 41,284 kg CO₂/ha-t bocsát ki. Ennek eredményeként a drónos technológia megközelítőleg 65 %-kal csökkenti a környezeti terhelést, ami jelentős előnyt jelent a fenntartható mezőgazdasági termelésben. A 65%-os CO₂-csökkenés és a 2,4-szeres energiamegtakarítás számításból származó kerekített arányok, ezért csak megközelítőleges értékek, nem közvetlenül a cikkben szereplő százalékok.¹²

García-Munguía és társai megállapították, hogy az AI-alapú drónrendszerek lehetővé teszik a precíziós permetezést akár 30%-kal kevesebb vegyszer felhasználása mellett. Emellett a permetezés pontossága fenntartható és akár tovább javítható is, ami csökkenti a környezeti terhelést.¹³

Nawaz és társai vizsgálatai szerint a traktorok és más mezőgazdasági gépek által okozott talajtömörödés jelentős hatással van a talaj fizikai tulajdonságaira, például a sűrűsége, ami megnehezíti a gyökerek növekedését és a tápanyagfelvételt. Ennek következtében a kukorica hozama 7,6% és 15,5% között csökkent azokon a területeken, ahol a talaj tömörödött, összehasonlítva a nem tömörödött területekkel. Bár Nawaz és társai vizsgálata nem hasonlította össze közvetlenül a drónos és a hagyományos permetezést, az eredmények alapján arra következtethetünk, hogy a talajt érintő fizikai behatások elkerülésével, például drónos permetezéssel a hozam csökkenése mérsékelhető.¹⁴

A dróntechnológia megjelenése forradalmasította a dokumentumfilm készítést azáltal, hogy lehetővé tette a korábban kizárólag nagy költségvetésű produkciók számára elérhető légi felvételek olcsó és rugalmas felhasználását, fokozatosan megszüntetve ezzel a légi felvételek luxus jellegét.¹⁵

7. Biztonságos leszállás az MI alapú VTOL segítségével

A tanulmány címe egyértelműen jelzi, hogy a vizsgálat középpontjában a mesterséges intelligencia a levegőben témaköre áll, vagyis annak elemzése, hogy a mesterséges intelligencia miként járult hozzá a dróntechnológia fejlődéséhez, valamint az iparág biztonságának és megbízhatóságának növeléséhez. A drónrendszerek számos képességgel rendelkeznek, amelyek teljes körű és részletes bemutatása meghaladná a jelen tanulmány kereteit. Ennek megfelelően a tanulmány nem törekszik valamennyi funkció ismertetésére, hanem a legfontosabb, a vizsgálat szempontjából kiemelt területekre koncentrálni.

A VTOL vagyis a Vertical Take-Off and Landing technológia lehetővé teszi, hogy a légi járművek függőlegesen emelkedjenek a levegőbe és ugyanilyen módon hajtsák végre a leszállást, ezáltal

¹² Safaeinejad et al., 2025

¹³ García-Munguía et al., 2024

¹⁴ Nawaz et al., 2023

¹⁵ Özgen, 2020

elkerülhető a hagyományos kifutópálya igénye. Ez a megoldás különösen hasznos olyan környezetben, ahol a rendelkezésre álló tér korlátozott, például sűrűn beépített városi területeken vagy nehezen hozzáférhető helyszíneken. A stabil és biztonságos működést korszerű szenzorrendszerek és intelligens irányítási algoritmusok támogatják. Ezek a rendszerek valós időben figyelik és folyamatosan korrigálják a jármű térbeli helyzetét, biztosítva az egyensúly és az irányíthatóság fenntartását. Az MI-alapú VTOL képesség képes a szenzoradatok valós idejű feldolgozására, önálló döntéshozatalra és a környezeti hatásokra történő gyors reagálásra.

A VTOL technológia létrejött nem a mesterséges intelligencia megjelenéséhez vagy fejlődéséhez köthető. Fontos kiemelni, hogy a függőleges fel- és leszállás lehetősége már az MI megjelenését megelőzően is létező és alkalmazott műszaki megoldás volt. Qin¹⁶ tanulmánya szerint a helikopterek voltak az első széles körben használt VTOL képességgel rendelkező repülőgépek, amelyeket már a II. világháborúban alkalmaztak. A VTOL repülőgépek fő kategóriái négy csoportba sorolhatók. Az első kategória az elektromos VTOL (eVTOL), amelyek elektromos hajtásúak, és akár ember szállítására is alkalmasak. A második kategória a hibrid VTOL, amely kombinálja a rotorokat és a hagyományos repülőgép-szárnyakat, így lehetővé téve a függőleges felszállást és leszállást, valamint a gyors vízszintes repülést. A harmadik kategória a rotoros helikopterek, amelyek nagy rotorlapátokkal biztosítják a függőleges felhajtóerőt, és precíz leszállásra is alkalmasak. A negyedik kategória a VTOL UAV-k vagy drónok, amelyek pilóta nélküli járművek különböző rotoros kialakításokkal, és főként fotózásra, videózásra használják őket.¹⁷

8. Összefoglaló

A tanulmány egyik legfontosabb eredménye, hogy átfogó képet ad a mesterséges intelligencia és a VTOL-technológia együttes szerepéről. A publikáció egyedisége abban rejlik, hogy az autonóm drónrendszereket nem kizárólag technológiai oldalról vizsgálja, hanem azok szervezeti és szakértői támogatási vonatkozásaira is kitér. A tanulmány rámutat, hogy az autonóm drónrendszerek nagyban segítik a precíziós permetezést és lehetővé teszik a megfizethető, minőségi felvételek készítését, ugyanakkor az emberi felügyelet, a szakértelem és a megfelelő szakértői támogatás továbbra is elengedhetetlen a biztonságos és hatékony működéshez. Az MI-alapú VTOL-technológia révén a nagyobb mezőgazdasági drónok biztonságosabbá és megbízhatóbbá váltak, ami a piaci kereslet növekedését is eredményezte. A rendelkezésre álló kutatások alapján az AI alkalmazása révén a permetezés akár 30%-kal kevesebb vegyszerrel végezhető és a CO₂-kibocsátás 65%-kal csökkenthető a hagyományos traktoros permetezéshez képest. Emellett a hagyományos mezőgazdasági műveletek során fellépő fizikai talajterhelés akár 5% feletti termésnövekedést is eredményezhet, amely drónos permetezéssel mérsékelhető.

Irodalomjegyzék

Caballero-Martín, D., Lopez-Guede, J.&M., Estevez, J. & Graña, M. (2024): *Artificial Intelligence Applied to Drone Control: A State of the Art*. Drones, vol. 8, no. 7, p. 296, DOI: 10.3390/drones8070296

Emmert-Streib, F., Yli-Harja, O. & Dehmer, M. (2020): *Artificial Intelligence: A Clarification of Misconceptions, Myths and Desired Status*. Frontiers in Artificial Intelligence, vol. 3, p. 524339, DOI:10.3389/frai.2020.524339

García-Munguía, A., Guerra-Ávila, P. L., Islas-Ojeda, E., Flores-Sánchez, J. L., Vázquez-Martínez, O., Margarito García-Munguía, A., & García-Munguía, O. (2024): *A review of drone technology and*

¹⁶ Qin, 2023

¹⁷ Qin, 2023

operation processes in agricultural crop spraying. *Drones*, 8(11), 674. <https://doi.org/10.3390/drones8110674>

Guebsi, R., Mami, S. & Chokmani, K. (2024): *Drones in Precision Agriculture: A Comprehensive Review of Applications, Technologies, and Challenges*. *Drones*, vol. 8, no. 11, p.686, DOI: 10.3390/drones8110686

Hassani, H., Silva, E. S., Unger, S., TajMazinani, M. & Mac Feely, S. (2020): *Artificial Intelligence (AI) or Intelligence Augmentation (IA): What Is the Future?*. *AI*, vol. 1, no. 2, pp.143–155, DOI:10.3390/ai1020008

Nawaz, M.M., Noor, M.A., Latifmanesh, H., Wang, X., Ma, W. & Zhang, W. (2023): *Field traffic-induced soil compaction under moderate machine-field conditions affects soil properties and maize yield on sandy loam soil*. *Frontiers in Plant Science*, vol. 14, p. 1002943, DOI: 10.3389/fpls.2023.1002943

Poór, J., Dajnoki, K., Kun, A. I., Pató G. Szűcs, B., Tóth, A., Szabó-Szentgróti, G., Kömüves, Z. S., Szabó, S. & Kálmán, B. G. (2025): *Trendek, tendenciák és paradigmaváltások az emberierőforrások menedzselése területén*. *Új Munkaügyi Szemle*, vol. 6, no. 3, pp. 2–13, DOI: 10.58269/umsz.2025.3.1

Qin, D. X. (2023): *Vertical takeoff and landing aircraft: Categories, applications, and technology*. *Proceedings of the 3rd International Conference on Computing Innovation and Applied Physics*, vol. 13, p. 124–129, DOI: 10.54254/2753-8818/13/20240810

Rejeb, A., Abdollahi, A., Rejeb, K., & Treiblmaier, H. (2022): *Drones in agriculture: A review and bibliometric analysis*. *Computers and Electronics in Agriculture*, vol. 198, p.107017, DOI: 10.1016/j.compag.2022.107017

Safaeinejad, M., Ghasemi-Nejad-Raeini, M., & Taki, M. (2025): *Reducing Energy and Environmental Footprint in Agriculture: A Study on Drone Spraying vs. Conventional Methods*. *PLOS ONE*, vol. 20, no. 6, e0323779, DOI: 10.1371/journal.pone.0323779

Talaeizadeh, A., Sharifi, I., Alasty, A. & Ghatreh Samani, S. (2025): *Agricultural Spraying Drones: A Comprehensive Review*. *Smart Agricultural Technology*, vol. 12, p. 101519, DOI:10.1016/j.atech.2025.101519

Tsouros, D. C., Bibi, S. & Sarigiannidis, P. (2019): *A Review on UAV-Based Applications for Precision Agriculture*. *Information*, vol. 10, no. 11, p. 349, DOI:10.3390/info10110349

Vlaskality, S. D., Lencsés, E. & Zalai, M. (2025): *Mezőgazdasági drónok alkalmazásának lehetőségei magyar szakértők véleményének feltérképezésével*. *GAZDÁLKODÁS: Scientific Journal on Agricultural Economics*, vol. 69, no. 01, pp. 35–49, DOI:10.22004/ag.econ.369149

Özgen, S. (2020): *The impact of drones in documentary filmmaking: Renaissance of aerial shot*. *AVANCA | CINEMA*, pp. 559-563 DOI:10.37390/ac.v0i0.74

Zhang, Z., Boubin, J., Stewart, C. & Khanal, S. (2020): *Whole-Field Reinforcement Learning: A Fully Autonomous Aerial Scouting Method for Precision Agriculture*. *Sensors*, vol.20, no.22, p.6585, DOI: 10.3390/s20226585

Felépíteni az ismeretlent – Autonóm, kooperatív jövőgépek és digitális ikrek a fenntartható földön kívüli infrastruktúra-fejlesztésben

Building the Unknown – Autonomous, Cooperative Future Machines and Digital Twins in Sustainable Extraterrestrial Infrastructure Development

Dr.¹ Ország-Krisz Axel² 

Absztrakt

A földön kívüli égitesteken, különösen a Holdon és a Marson történő tartós emberi jelenlét és a hosszú távú űrkutatási missziók fenntarthatóságának alapvető záloga az in-situ erőforrás-hasznosítás (ISRU, In-Situ Resource Utilization). Jelen tanulmány egy olyan transzdiszciplináris módszertani keretrendszerrel mutat be, amely a digitális iker (Digital Twin) technológia, a tanterv-alapú gépi tanulás (Curriculum Learning) és a sztochasztikus bizonytalanság-kezeléssel felruházott kooperatív robotika integrációjára épül. Célunk egy olyan 7 lépéses rendszerszintű folyamatmodell prezentálása, amely képes áthidalni a földi analóg környezetek és a determinisztikusan nem modellezhető űrbéli fizikai valóság közötti transzfer-rést (Sim-to-Real gap). A javasolt modellt egy holdbéli menedék- és leszállóhely autonóm földmunkáinak és infrastrukturális kiépítésének esettanulmányán keresztül validáljuk, kritikai elemzés alá vetve a jelenlegi nemzetközi élvonalbeli technológiákat (NASA RASSOR, DLR ARCHES). A tanulmány bizonyítja, hogy az autonómia és a szigorú szabványkövetés szimbiózisa, valamint a heterogén robotflották dinamikus együttműködése a földön kívüli mérnöki tevékenység megkerülhetetlen paradigmája.

Abstract

In-Situ Resource Utilization (ISRU) is a fundamental key to the sustainability of long-term human presence and space exploration missions on extraterrestrial bodies, especially on the Moon and Mars. This paper presents a transdisciplinary methodological framework that integrates Digital Twin technology, Curriculum Learning, and cooperative robotics with stochastic uncertainty management. Our goal is to present a 7-step system-level process model that can bridge the Sim-to-Real gap between terrestrial analog environments and the physical reality in space that cannot be deterministically modeled. The proposed model is validated through a case study of autonomous groundworks and infrastructure construction of a lunar shelter and landing site, critically analyzing current international cutting-edge technologies (NASA RASSOR, DLR ARCHES). The study demonstrates that the symbiosis of autonomy and strict adherence to standards, as well as the dynamic cooperation of heterogeneous robot fleets, is an unavoidable paradigm for extraterrestrial engineering.

¹ Jogász doktor (nem tudományos fokozat)

² mesterséges intelligencia fejlesztő és szakértő, adattudós, email@drorszagkriszaxel.hu

1. Bevezetés

1.1. Az in-situ erőforrás-hasznosítás (ISRU) gazdaságtana és logisztikája

A 21. század első felének űrkutatását, különösen a NASA Artemis-programját és a nemzetközi partnerintézmények Hold- és Mars-misszióit egy alapvető paradigmaváltás jellemzi: a Föld-centrikus, merev ellátási láncokra épülő logisztikai modellek végérvényesen elérték fenntarthatósági korlátaikat. A korábbi Apollo-korszak missziós architektúrája arra a koncepcióra épült, hogy az életfenntartáshoz, az infrastrukturális elemekhez, valamint a visszatéréshez szükséges minden egyes gramm anyagot, beleértve a strukturális elemeket, a vizet, az oxigént és a hajtóanyagot, a Föld gravitációs kútjából kell kiszabadítani és a célállomásra szállítani.

Ez a megközelítés lineárisan és fenntarthatatlanul növeli a missziók költségeit a távolság és a tartózkodási idő növekedésével. Ezzel szemben az In-Situ Resource Utilization (ISRU), azaz a helyszíni erőforrás-hasznosítás elve kimondja, hogy a földön kívüli égitesteken található nyersanyagokat (regolitot, illékony anyagokat, jéglerakódásokat) helyben kell kitermelni, feldolgozni és strukturális vagy energetikai célokra felhasználni. Az ISRU nem csupán egy technológiai opció, hanem a tartós bolygóközi jelenlét gazdasági sine qua nonja (elengedhetetlen feltétele).

A helyszíni anyagfelhasználás révén a missziók kezdeti tömege (IMLEO: Initial Mass in Low Earth Orbit) radikálisan csökkenthető. Ha az építkezésekhez szükséges strukturális védőelemeket, például a kozmikus sugárzás és a mikrometeoritok elleni regolit-pajzsokat, nem a Földről indítjuk, a hasznos teher aránya a tudományos berendezések javára tolódik el. Ezen túlmenően az ISRU biztosítja az azonnali reakcióképesség előnyét: a misszió során felmerülő váratlan infrastrukturális igények (pl. újabb sugárvédelmi sáncok, hőtároló tömegek, vagy leszállóhely-károsodások korrekciója) azonnal, helyszíni nyersanyag-utánpótlásból fedezhetők, függetlenül a földi indítási ablakoktól és a több hónapos transzferidőktől.

1.2. A Föld-centrikus ellátási láncok fizikai korlátai

Az ISRU szükségességét legplasztikusabban a klasszikus rakétaegyenlet, a Ciolkovszkij-egyenlet határozza meg, amely bemutatja a hordozórakéták tömegarányainak exponenciális növekedését a szükséges sebességváltozás (Δv) függvényében:

$$v(t) = v_g \cdot \ln\left(\frac{m_0}{m(t)}\right)$$

Ahol:

v a rakéta sebessége a t időpillanatban,

v_g a rakétát elhagyó gáz sugár sebessége a rakétához képest (jellemző érték kémiai hajtóanyag esetén: 4,5 km/s),

m_0 a rakéta induló tömege és

m a rakéta tömege az indulástól számított t idő múlva.

1. ábra: Ciolkovszkij-egyenlet

Ha egy holdi vagy marsi bázis felépítéséhez szükséges nehéz szerkezeti elemeket, talajmozgató gépeket és azok ballasztanyagait a Földről kívánjuk eljuttatni a célállomás felszínére, a láncolt Δv igények (Alacsony Föld körüli pálya $\rightarrow$ Holdirányú pálya $\rightarrow$ Hold körüli pálya $\rightarrow$ Holdfelszín) miatt az m_0 / m_f arány drasztikusan megemelkedik. Minden egyes kilogrammnyi hasznos teher

felszínre juttatása a Holdon a hordozórakéta indítási tömegét tekintve nagyságrendileg 15-20 kilogrammnyi plusz tömeget igényel a Földön. A Mars esetében ez az arány még kedvezőtlenebb.

Ebből a fizikai axiómából következik, hogy a földön kívüli nehéz-infrastruktúrát „helyben kell megtermelni”. Nem a kész épületeket vagy a nehéz betonstruktúrákat kell az űrbe juttatni, hanem azokat a magas intelligenciával és adaptivitással rendelkező „jövőgépeket”, autonóm robotrendszereket, amelyek képesek a helyszínen rendelkezésre álló szilikátokból, fémoxidokból és regolitból felépíteni a szükséges létesítményeket. Ezzel a logisztikai lánc fókuszpontja a tömegszállításról az információsűrűsége és az algoritmikus autonómiára helyeződik át.

2. A Földön Kívüli Környezet Determinisztikus Modellezésének Korlátai

2.1. Regolit-dinamika: koptató hatás, elektrosztatikus feltöltődés és vákuum-kohézió

A földi bányászatban, mélyépítésben és automatizálásban alkalmazott mechanikai modellek közvetlenül nem transzponálhatók a földön kívüli környezetre, mivel az ott jelen lévő fizikai tényezők komplex és sokszor nemlineáris interakcióban állnak a gépi rendszerekkel. A legjelentősebb kihívást a holdi és marsi felszínt borító regolit jelenti.

A földi talajokkal ellentétben, ahol a szél- és vízerózió koptatja, gömbölyíti a szemcséket, a holdi regolitot évmilliárdos mikrometeorit-bombázás hozta létre. Ennek következtében a regolitzemcsék rendkívül élesek, szögletesek és üvegszerűen abrázivadnak (koptató hatásúak). Amikor egy autonóm munkagép talajmozgatási vagy kotrési tevékenységet végez, a mechanikai alkatrészek (csapágyak, lánctalpak, hidraulikus tömítések) olyan mértékű tribológiai igénybevételnek vannak kitéve, amely a földi üzemidők töredéke alatt teljes mechanikai meghibásodáshoz, a tömítések átszakadásához és a mozgó alkatrészek berágódásához vezet.

A légkör hiánya (vákuumkörnyezet) és a direkt napsugárzás (UV és röntgen) egy másik kritikus jelenséget generál: az elektrosztatikus feltöltődést. A holdi nappali oldalon a fotoelektromos hatás miatt a regolit pozitív töltésűvé válik, míg az éjszakai oldalon és a permanensen árnyékos régiókban (PSR, Permanently Shadowed Regions) a plazmakörnyezet negatív töltést indukál. Ez a finomszemcsés port (PM2.5 alatti frakciók) lebegtetésre és a robotok felületéhez való makacs tapadásra készíti. A tapadó porréteg drasztikusan csökkenti a napelemek hatásfokát, gátolja a hőszugárzók működését (termikus megfutást okozva), és szoftveresen torzítja az optikai szenzorok (LiDAR, kamerák) jeleit.

Ugyancsak a nagyvákuum környezet számlájára írható a vákuum-kohézió jelensége. A gázmolekulák hiánya miatt a regolitzemcsék felületéről hiányzik az az adszorbeált gázréteg, amely a Földön természetes kenőanyagként működik. Ennek következtében a tiszta ásványi felületek között fellép a hideghegedés (cold welding) és a megnövekedett belső súrlódás (kohézió). A talaj nyírószilárdsága és viselkedése a mélységgel nemlineárisan változik, amit a klasszikus Terzaghi-féle talajmechanikai egyenletekkel nem lehet megbízhatóan előre jelezni.

A gravitáció radikális csökkenése megváltoztatja a gépek tapadását és reakcióerejét is. Egy földi kotrógép a saját súlyából adódó tapadási súrlódást használja fel a kanál földbe mélyesztéséhez. A Holdon egy ugyanilyen tömegű gépnek hatszor kisebb a súlya, így a kanál benyomásakor fellépő reakcióerő egyszerűen felemeli a gép elejét a talajról, mielőtt a regolit elnyíródna. Emiatt teljesen új geometriájú és kinematikájú eszközök (pl. ellenirányban forgó ikerdobos marófejek, mint a NASA RASSOR esetében) tervezése és vezérlése szükséges.

2.2. A klasszikus hibakorrekció és a humán intervenció hiányának rendszerszintű kockázatai

A földi automatizált rendszerek és az Ipar 4.0 struktúrák mögött mindig ott áll a közvetlen humán intervenció lehetősége. Ha egy autonóm raktári robot vagy egy bányászati célgép szoftveres vagy

mechanikai anomáliát észlel, a rendszer vészleállást hajt végre, és a karbantartó személyzet órákon belül képes elhárítani a hibát (alkatrészcsere, fizikai újraindítás, szoftveres hibakeresés).

A földön kívüli infrastruktúra-építés fázisában ez a biztonsági háló nem létezik. A Föld és a Hold közötti oda-vissza kommunikációs késleltetés (latencia) átlagosan ~2,56 másodperc, míg a Mars esetében ez az orbitális pozícióktól függően 6 és 44 perc között változik. Ez a távolság kizárja a valós idejű teleoperációt (távvezérlést). Ha egy robotgép a Holdon egy instabil árokparton megcsúszik, a földi operátor mire észleli a telemetriából a dőlést és kiküldi a „Stop” parancsot, a gép már felborult.

Ebből adódóan a rendszereknek teljes körű kognitív autonómiával kell rendelkezniük. A klasszikus determinisztikus programozás, amely előre rögzített elágazásokkal (If-Then-Else struktúrákkal) kezeli a környezeti változókat, képtelen lefedni az idegen égitesteken fellépő váratlan szituációk végtelen kombinációját (pl. egy váratlanul felbukkanó, a regolit alatt rejtőző bazalttömb, a lánc talpba szoruló éles kőzetdarab, vagy a sugárzás okozta bit-átfordulás [single-event upset] a fedélzeti számítógép memóriájában). A hibakorrekció hiánya miatt minden szoftveres és fizikai döntésnek „elsőre jónak” kell lennie, mivel egyetlen kritikus hiba az egész misszió elvesztését, és ezáltal dollármilliárdos kárt jelenthet.

3. Elméleti Módszertan: Digitális Iker és Kockázatkezelő Mesterséges Intelligencia

3.1. A Digitális Iker architektúrája és valós idejű szinkronizációja

A földön kívüli autonóm építkezések tervezési és végrehajtási fázisa közötti szakadék áthidalására a legígéretesebb technológiai megközelítés a kétrétegű Digitális Iker (Digital Twin, DT) architektúra implementálása. Az űrkutatási kontextusban a Digitális Iker nem csupán egy statikus CAD-modell vagy egy egyszerűsített vizualizációs felület, hanem egy olyan dinamikus, többszintű kiber-fizikai rendszer (CPS), amely valós időben, a kommunikációs késleltetés (latencia) korlátain belül, szinkronizálja a célobszektor fizikai állapotát annak digitális reprezentációjával.

A rendszer architektúrája két fő komponensre bontható fel:

- **Mikro-szintű:** (Edge-alapú) Digitális Iker: Közvetlenül a földön kívüli robotok fedélzeti számítógépein (On-Board Computers, OBC) fut. Feladata a gép közvetlen fizikai interakcióinak leképezése, mint például a lánc talp és a regolit közötti csúszás (slip ratio), a kotrókar hidraulikus vagy elektromos nyomatékigénye, valamint a belső rendszerek termikus gradienseinek modellezése. Ez a réteg felelős a valós idejű, ezredmásodperces nagyságrendű reflexszerű korrekciókért.
- **Makro-szintű:** (Központi/Földi) Digitális Iker: A földi irányítóközpont nagyteljesítményű számítástechnikai klaszterein üzemel. Ez magában foglalja a teljes holdbéli vagy marsi munkaterület háromdimenziós, nagy felbontású topográfiai modelljét, amelyet a felderítő roverek és keringőegységek LiDAR, valamint fotogrammetriai adatai folyamatosan frissítenek. Itt futnak a több ágenses flottaoptimalizációs algoritmusok és a hosszú távú trajektóriatervezők.

A két szint közötti szinkronizáció a sávszélesség-korlátok és a terjedési késleltetés miatt transzformált állapotter-átvitellel valósítható meg. A fizikai robot nem a teljes nyers szenzoros adatfolyamot küldi vissza a Földre, hanem egy tömörített, matematikai varianciákkal leírt állapotvektort. A földi Makro-Iker ebből az állapotvektorból, valamint a diszkrét időközönként küldött háromdimenziós pontfelhőkből (Point Cloud) rekonstruálja a környezet változásait.

Ha a fizikai valóság és a szimulált modell között eltérés (divergencia) lép fel, például a regolit belső súrlódási szöge a valóságban magasabb, mint amit a földi laboratóriumi analógok alapján becsültek, a Makro-Iker automatikusan újralibrlja a szimulációs motor paramétereit a

Kalman-szűrők (Extended Kalman Filter, EKF) egy adaptív variánsával, majd a frissített viselkedési szabályokat visszasugározza az űrbéli eszközöknek.

3.2. Curriculum Learning (Tanterv-alapú tanulás) a Markov-féle döntési folyamatokban (MDP)

A robotok autonóm döntéshozatali mechanizmusait és mozgástervezését megerősítéses tanuláson (Reinforcement Learning, RL) alapuló modellek vezérlik. Egy olyan komplex környezetben, mint egy idegen égitest felszíne, a robot cselekvései egy kontinuus állapottérrel rendelkező Markov-féle döntési folyamatként (Markov Decision Process, MDP) definiálható, amelyet az S, A, P, R, γ ötessel írunk le, ahol:

- S az állapotok halmaza,
- A az akciók halmaza,
- P az állapotátmeneti valószínűségegyüttható,
- R a jutalomfüggvény (Reward Function),
- γ a diszkonttényező.

Tekintettel arra, hogy a véletlenszerű próbálkozásokon alapuló klasszikus RL algoritmusok (pl. PPO, SAC) egy holdbéli bázisépítés komplexitása mellett a jutalmak ritkasága (sparse rewards) miatt konvergencia-hibával vagy végtelen tanulási idővel szembesülnének, a rendszerben Tanterv-alapú tanulást (Curriculum Learning, CL) is érdemes alkalmazni. A CL lényege, hogy a tanulási folyamatot diszkrét, egymásra épülő nehézségi szintekre (kurzusokra) bontja:

Ahol minden egyes szint egy-egy specifikus feladat-kontextust és környezeti eloszlást reprezentál. A holdbéli talajszintezési feladatra lefordítva a tanterv felépítése például a következő lehet:

1. Alapozó kurzus: Kinematikai stabilitás fenntartása sík, súrlódásmentes felületen, akadályok nélkül. A jutalomfüggvény kizárólag a célpozíció elérését és az orientáció megőrzését díjazza.
2. Szenzoros és morfológiai kurzus: Navigáció és mozgás egyenetlen terepen, fix méretű kráterek és sziklák elkerülésével, idealizált kohéziójú regolitban.
3. Interakciós kurzus: A kotró-maró szerszámfej aktiválása. A robot megtanulja adagolni a vágási erőt a gép dőlésszögének és a lánctalp tapadásának függvényében.
4. Komplex kooperatív kurzus: Anyagmozgatás dinamikusan változó topográfia mellett, ahol a környezetet más robotok tevékenysége és a folyamatosan képződő pordepóniák egyidejűleg módosítják.

A tanulás átvitele az egyes fázisok között a neurális hálózat súlymatricáinak transzferével történik. Ez a megközelítés drasztikusan csökkenti a szimulációs időigényt, és minimalizálja az úgynevezett „katasztrofális felejtés” (catastrophic forgetting) kockázatát, biztosítva, hogy a robot stabil alapreflexekkel rendelkezzen a kritikus fizikai műveletek során.

3.3. Sztochasztikus ML: Episztemikus és aleatórikus bizonytalanságok számszerűsítése (Bayesi neurális hálók)

Az autonóm űreszközök biztonságos működéséhez elengedhetetlen, hogy a fedélzeti Mesterséges Intelligencia ne csupán döntéseket hozzon, hanem képes legyen saját döntéseinek megbízhatóságát és a bejövő adatok bizonytalanságát is precízen kvantifikálni. A kockázatkezelési modellekben szigorúan kettéválasztható a bizonytalanság két alaptípusa:

1. Aleatórikus bizonytalanság (Adatbizonytalanság): A fizikai mérőrendszerek belső zajából, a szenzorok felbontási korlátaiból és a környezeti sztochasztikus hatásokból (pl. a porszemcsék okozta lézersugár-szóródás a LiDAR szenzorban) adódik. Ez a

bizonytalanság nem csökkenthető több adat gyűjtésével, a modellnek közvetlenül meg kell tanulnia a bemeneti adatok varianciáját.

2. **Episztemikus bizonytalanság (Modellbizonytalanság):** Abból fakad, hogy a neurális hálózat képzési adathalmaza nem fedte le tökéletesen a valós működési teret (out-of-distribution adatok). Ez a bizonytalanság elméletileg csökkenthető lenne további szimulációkkal vagy mérésekkel.

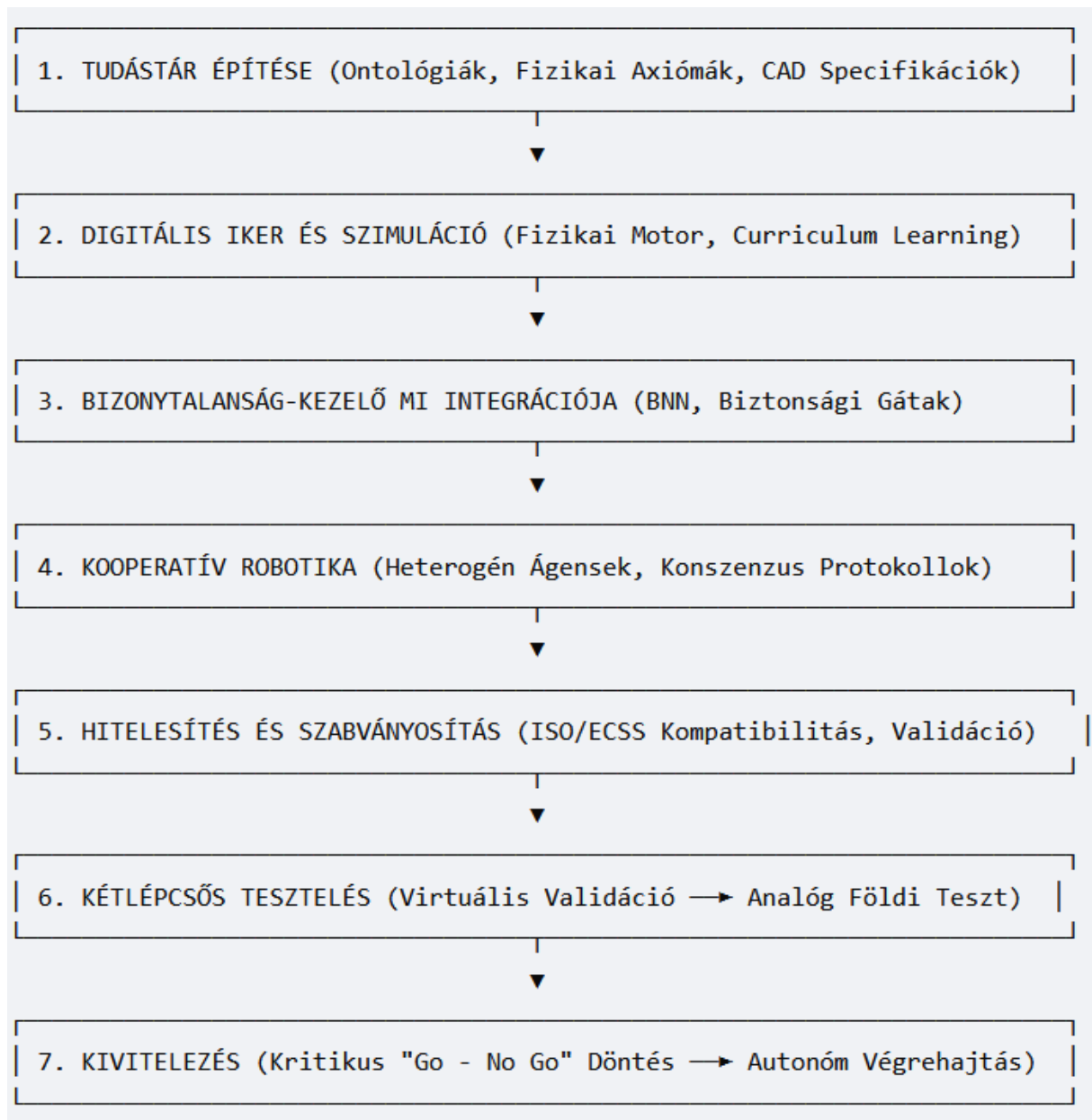
A bizonytalanságok számszerűsítésére Bayesi Neurális Hálózatok (Bayesian Neural Networks, BNN) és Monte Carlo Dropout technikák alkalmazhatók. A hagyományos neurális hálózatok determinisztikus súlyai helyett a BNN a súlyok feletti valószínűségi eloszlást határoz meg, amely egy poszteriór eloszlással írható le a D tanító adatbázis ismeretében. Mivel a poszteriór eloszlás egzakt kiszámítása analitikusan intolerábilis komplexitású, variációs következtetés (Variational Inference) alkalmazható annak érdekében, hogy minimalizálható legyen a Kullback-Leibler (KL) divergencia egy egyszerűbben kezelhető eloszláscsalád és a valódi poszteriór között:

Inferencia közben a robot OBC-je minden döntési ciklusban M számú sztochasztikus előrecsatolást (forward pass) hajt végre a hálózaton aktív dropout mellett. Az így kapott predikciók eloszlásának empirikus varianciája közvetlenül megadja az episztemikus bizonytalanságot.

Ha egy munkagép olyan talajszerkezettel vagy geometriai formációval találkozik, amely a tudástárában és a szimulációs előtörténetében nem szerepelt, az episztemikus bizonytalanság értéke átlépi a kritikus biztonsági küszöbértéket. Ebben a pillanatban az MI aktiválja a biztonsági gátat: leállítja a fizikai végrehajtást, rögzíti a gép aktuális pozícióját, és távsegítséget (High-Level Tele-Assistance) kér a földi irányítóközponttól, elkerülve a gép önveszélyes viselkedését.

4. A Földön Kívüli Infrastruktúra-Építés 7 Lépéses Folyamatmodellje

A bizonytalanságok minimalizálására és a strukturális integritás garantálására egy szigorúan szekvenciális, de visszacsatolási hurkokkal ellátott 7 lépéses rendszerszintű folyamatmodellt állítottunk fel. Ez a keretrendszer biztosítja a technológiai transzfert a tiszta elméleti tudománytól a fizikai megvalósításig.



2. ábra: A 7 lépéses modell

1. Lépés: *Ontológiai és tudományos tudástár felépítése*

A folyamat egy transzdiszciplináris adatbázis és formális tudásreprezentáció (ontológia) létrehozásával kezdődik. A tudástár három fő pillérre tagozódik:

- Tudományos alapismeretek: Geokémiai és ásványtani adatok (regolit összetétele, ilmenit-, anortit-tartalom), termodinamikai egyenletek (szélsőséges hőmérsékleti profilok), radiológiai modellek (galaktikus kozmikus sugárzás és szoláris részecskeesemények fluxusa), valamint mikrogravitációs pordinamika.
- Mérnöki-kivitelezési tartomány: Robotikai kinematikai modellek, bányászati és jövesztési technikák adaptációi, additív gyártási (3D nyomtatási) paraméterek (lézeres szinterezés energiasűrűsége, mikrohullámú olvasztási frekvenciák) és talajmechanikai határadatok.
- Berendezés-specifikus adatbázisok: A misszióban részt vevő összes fizikai eszköz részletes specifikációja, beleértve a CAD-modelleket, a motorok hatásfokgörbéit, a szenzorok zajkarakterisztikáit és a redundáns alrendszerek hibafáit (Fault Tree Analysis).

2. Lépés: Digitális iker generálása és szimulációs környezet építése

A tudástár adataira támaszkodva felépül a nagypontosságú, determinisztikus és sztochasztikus elemeket ötvöző szimulációs környezet. A szimulációs motorban a Diszkrét Elem Módszer (Discrete Element Method, DEM) alkalmazása ajánlott a regolitszemcsék egymásra hatásának modellezésére, ami elengedhetetlen a kotrási és tömörítési erők valóságghű szimulációjához. Ebben a fázisban indul el a robotok irányító szoftvereinek Curriculum Learning alapú tömeges párhuzamosítása (Massively Parallel Simulation), virtuálisan több ezer üzemórát generálva másodpercek alatt. Hangsúlyozandó, hogy ez a lépés teljes egészében a felbocsátás előtt, földi számítási kapacitásokon realizálódik.

3. Lépés: A bizonytalanságra tanított és biztosított MI integrációja

A szimulációs fázisban kiképzett neurális hálózatokat fel kell ruházni a 3.3. alfejezetben bemutatott sztochasztikus bizonytalanság-quantifikáló rétegekkel. Az MI ágensekbe merev matematikai korlátokat (Safe Reinforcement Learning) és determinisztikus biztonsági logikai kapukat kell integrálni. Ezen gátak feladata megakadályozni, hogy a tanuló algoritmus olyan extrém akciókat hajtson végre (pl. túl meredek dőlésszögű instabil haladás), amelyek a hardver azonnali pusztulásához vezetnének.

4. Lépés: Kooperatív robotikai protokollok konfigurálása

Mivel a földön kívüli infrastruktúra-építés meghaladja egyetlen monolitikus gép képességeit, ebben a fázisban történik a heterogén robotflotta közötti kommunikációs és feladatmegosztási protokollok implementálása. Az ágensek közötti koordinációt egy elosztott, decentralizált konszenzus-algoritmus (pl. piac-alapú aukciós protokollok vagy Distributed Constraint Optimization, DCOP) vezérli. A robotok autonóm módon licitálnak az egyes részfeladatokra (pl. felderítés, szállítás, tömörítés) a saját aktuális energiaszintjük, pozíciójuk és mechanikai állapotuk alapján, minimalizálva a teljes flotta energiafelhasználását és kiküszöbölve az egyponos rendszerhibákat (Single Point of Failure).

5. Lépés: Hitelesítés, szabványosítás és ellenőrzés

A szoftveres és hardveres architektúrát alá kell vetni az űripari szabványok szigorú ellenőrzési folyamatának. Ez magában foglalja az Európai Űrutazási Szabványosítási Együttműködés (ECSS: European Cooperation for Space Standardization) szoftver-megbízhatósági (ECSS-E-ST-40C) és minőségbiztosítási (ECSS-Q-ST-80C) előírásainak való megfelelést. A neurális hálózatok fekete doboz jellege formális verifikációs módszerekkel (pl. Abstract Interpretation vagy Satisfiability Modulo Theories, SMT alapú hálózati verifikáció) kompenzálható annak igazolása érdekében, hogy a gép irányítófüggvénye a megengedett állapotter belső határain belül marad.

6. Lépés: Kétlépcsős tesztelés és validáció

A fizikai megvalósítás előtti végső minősítési fázis:

- **Szoftver-a-hurokban (Software-in-the-Loop, SIL) és Hardver-a-hurokban (Hardware-in-the-Loop, HIL) tesztelés:** Az MI szoftvert át kell telepíteni a végső, sugárzástűrő célhardverre (pl. RAD750 vagy specifikus FPGA architektúrák), és a bemeneteket a Digitális Iker által generált szintetikus szenzorjelekkel kell táplálni, mérve a számítási késleltetést és a valós idejű végrehajtási stabilitást.
- **Fizikai analóg tesztelés:** A robotok földi ikerpéldányait (Engineering Models) kell elhelyezni egy fizikai analóg környezetben (pl. vulkáni kráterek vagy speciális regolit-szimulációs csarnokok). Itt ellenőrizhető a mechanikai struktúrák tényleges kopása, a

porbejutást és a flotta együttműködési hatékonysága valós fizikai körülmények és kényszerúségek mellett.

7. Lépés: "Go - No Go" döntéshozatal és autonóm kivitelezés

A misszió indítása és a célégitestre való megérkezés után a rendszer végrehajtja a hardveres önellenőrzést (Built-In Self-Test, BIST). Az összegyűjtött lokális telemetria és a földi Digitális Iker predikcióinak összevetése alapján a mérnökcsoport és a központi felügyelő MI közösen meghozza a kritikus „Go - No Go” döntést. Pozitív („Go”) döntés esetén a flotta megkezdheti a teljesen autonóm, helyszíni infrastrukturális kivitelezést, miközben a fedélzeti és földi Digitális Iker folyamatosan, zárt szabályozási körben monitorozza és dokumentálja a haladást.

5. Esettanulmány: Holdbéli Menedék- és Leszállóhely Létesítése

5.1. A probléma felbontása (Work Breakdown Structure - WBS) és geometriai kritériumok

A bemutatott elméleti módszertan és a 7 lépéses folyamatmodell gyakorlati verifikációját egy kritikus fontosságú holdbéli infrastruktúra-projekt, egy kombinált menedékhely (habitat védőburkolat) és egy rakéta-leszállóhely (landing pad) autonóm kivitelezési tervén keresztül mutatom be. A sugárzásvédelem és a mikrometeoritok elleni védelem érdekében a lakómodulokat egy minimum 1,5–2 méter vastagságú tömörített regolit-pajzzsal kell befedni. A leszállóhely esetében a legfőbb mérnöki kihívást a rakétahajtóművek gázsugarának koptató és talajdiszperziós hatása (plume-surface interaction) jelenti. Megfelelő talajelőkészítés hiányában a hajtóműből kiáramló szuperszonikus gázáram másodpercenként több száz kilogrammnyi regolitot gyorsít fel hipersebességre, amely csiszolóvászonnként semmisítheti meg a környező bázisstruktúrákat, kommunikációs antennákat és napelemparkokat.

A projekt rendszerszintű komplexitását egy szigorú Funkcionális Felbontási Struktúra (Work Breakdown Structure, WBS) mentén dekomponáljuk, amely meghatározza az egymásra épülő mérnöki fázisokat és a hozzájuk rendelt geometriai tűréshatárokat.

A leszállóhely geometriai kritériumai megkövetelik a teljes síklapúságot: a maximális megengedett dőlésszög a felület egyetlen pontján sem haladhatja meg az 1 %-ot, míg a finomszintezés után visszamaradó lokális felületi érdesség (surface roughness) is limitált. A talajtömörítési fázis célja, hogy a regolit sűrűségét a természetes alapértékről egy kritikus tartományba növelje. Ez biztosítja, hogy a talaj teherbíró képessége (California Bearing Ratio, CBR) meghaladja a 80 %-ot, minimalizálva a leszálló úrhajók talpának besüllyedését és a gázsugar eróziós kráterképződését.

5.2. Heterogén flotta-architektúra: Felderítők, nehézgépek és szuborbitális drónok kooperációja

Ezen összetett, egymással párhuzamosan futó feladatok végrehajtásához egy optimalizált, heterogén ágensekből álló robotflotta-architektúrát kell konfigurálni. A flotta tagjai funkcionálisan differenciáltak, azonban közös elosztott kommunikációs hálón (Ad-Hoc Mesh Network) keresztül osztják meg állapotér-információikat:

1. Autonóm Felderítő Egységek (Scout Rovers): Kis tömegű, nagy mobilitású roverek, amelyek multispektrális kamerákkal, sztereó-látórendszerekkel és talajba hatoló radarral (Ground Penetrating Radar, GPR) vannak felszerelve. Feladatuk a felderítési fázis végrehajtása: a felszín alatti anomáliák (pl. rejtett bazalttömbök, üregek) detektálása és a mikrotipológiai térkép átvitele a földi és fedélzeti Digitális Ikerbe.
2. Nehéz Munkagépek (Heavy Operational Machines):

- a. Kotró-maró egység (Excavator): Ellenirányban forgó ikerdobos marófejjel szerelt lánc talpas gép. Az ellenirányú forgatónyomaték kompenzálja a holdi mikrogravitáció (1/6 g) okozta tapadásvesztést, így a gép anélkül képes nagy mennyiségű regolit kitermelésére, hogy felemelkedne a talajról.
 - b. Szállítógép (Hauler/Dumper): Nagy befogadóképességű, rugalmas kerékfelfüggesztésű ágens, amely a kitermelt regolitot mozgatja a dömperzónák és a menedékhely építési területe között.
 - c. Földgyalu és Tömörítő gép (Grader & Compactor): Célrányos profilozó tolólappal és vibrációs tömörítő hengerrel ellátott nehézgép, amely az előkészítő fázisok fizikai kivitelezéséért felelős. A tömörítő fej energiahatékonyságát nagyfrekvenciás piezoelektromos oszcillátorok optimalizálják.
3. Ellenőrző Szuborbitális Drónok (Inspection Drones): Mivel a Holdon a légkör hiánya kizárja az aerodinamikai repülést, ezek az egységek miniatűr, hideggázos (Cold-Gas) impulzushajtóművekkel végrehajtott ballisztikus ugrásokkal változtatnak pozíciót. Feladatuk a munkaterület feletti periodikus átrepülés során készített nagy felbontású LiDAR pontfelhők segítségével a szintezési és tömörítési fázisok milliméter-pontos minőségbiztosítása (QA).

A flotta kooperációját a III. fejezetben említett piac-alapú aukciós protokoll vezérli. Amikor a szállítógép jelzi, hogy kapacitása elérte a 100%-ot, a rendszer automatikusan diszpécser-aukciót indít az optimális lerakóhelyre. Ha a tömörítő gép szenzorai azt észlelik, hogy egy adott zónában a talaj sűrűsége még nem érte el a célértéket, a zónát zárolja az additív építési fázis elől, és újraosztja a tömörítési részfeladatot a szabad kapacitással rendelkező nehézgépek között.

6. Globális Technológiai Körkép és TRL-Szintek Kritikai Elemzése

A bemutatott autonóm keretrendszer elméleti konstrukciói nem elszigetelt modellek; a nemzetközi űrkutatási ügynökségek (NASA, ESA, DLR) és azok akadémiai partnerei az elmúlt időszakban számos olyan hardveres és szimulációs komponenst fejlesztettek ki, amelyek közvetlen építőelemei a földön kívüli infrastruktúra-építésnek. Az alábbiakban kritikai elemzés alá vetjük ezen élvonalbeli technológiák Technológiai Készültségi Szintjét (Technology Readiness Level, TRL) és azok korlátait.

6.1. NASA RASSOR, DLR ARCHES és az Etna-analóg missziók tanulságai

A NASA Regolith Advanced Surface Systems Operations Robot (RASSOR) projektje kifejezetten a holdi mikrogravitációs környezet talajkitermelési problémájára adott mérnöki válasz. A RASSOR szakít a hagyományos földi kotrógépek felépítésével: két végén elhelyezett, ellenirányban forgó dobos marófejeket (bucket drums) alkalmaz. A TRL 6-os szinten álló rendszer prototípusai vákuumkamrás és csökkentett gravitációs repülések (parabolikus repülések) során bizonyították, hogy a vágási reakcióerők belső kiegyenlítése révén a gép képes saját tömegénél nagyobb mennyiségű regolit megmozgatására anélkül, hogy elveszítené tapadását.

Ugyanakkor a RASSOR korlátja, hogy könnyű szerkezete miatt képtelen a mélyebb, erősen tömörödött talajrétegek jövesztésére. Ezt a rést hivatott betölteni a NASA Kennedy Space Center által fejlesztett KSC-TOPS-7 nehéz-kotrógép prototípus, amely robusztusabb mechanikai felépítéssel és speciális rezgőkésekkel kísérli meg a mélyreható talajmunkák elvégzését. A TOPS-7 jelenleg TRL 4/5-ös szinten áll; a legnagyobb kihívást a hidraulikus rendszerek teljes kiváltása jelenti elektromechanikus és piezokeramikus aktuátorokkal, mivel a hidraulikaolajok viszkozitása a holdi kriogén éjszakák során összeomlást idézne elő.

6.2. Additív gyártástechnológiák regolit-alapon: Regolight és 3D Habitat Challenge

A robotok közötti kooperáció és az autonóm környezeti adaptáció terén a Német Űrkutatási Központ (DLR) ARCHES (Autonomous Robotic Networks to Help Modern Societies) elnevezésű világűr-analóg demonstrációs missziója mutatott be áttörést. A Catania melletti Etna vulkán lankáin végrehajtott terepi tesztek során heterogén robotflották (felderítő roverek és manipulátorkarral szerelt végrehajtó egységek) teljesen autonóm módon, humán beavatkozás nélkül hajtottak végre geológiai mintavételi és hálózati kiépítési feladatokat. Az ARCHES sikeresen verifikálta a decentralizált hálózati topológiák és az elosztott feladatallokáció működőképességét valós, instabil morfológiájú, koptató hatású vulkáni hamuval borított környezetben (TRL 7).

A szimulációs és fizikai tesztkörnyezetek integrációjának európai központja a Kölnben található DLR/ESA LUNA Facility. Ez a létesítmény egy 1000 négyzetméteres zárt csarnok, amelyet teljesen feltöltöttek EAC-1A típusú holdi regolit-szimulánssal. A LUNA csarnok egyedülálló tulajdonsága a finomhangolható, felfüggesztéses gravitáció-szimulációs rendszer, valamint a napfény-szimulátor, amely képes reprodukálni a holdi sarkvidékekre (South Pole) jellemző extrém alacsony beesési szögű, éles árnyékokat generáló megvilágítást. A LUNA környezet közvetlen fizikai interfészként szolgál a II. fejezetben tárgyalt Digitális Ikerk verifikációjához (HIL tesztkörnyezet, TRL 6).

6.3. Additív technológiák és strukturális konszolidáció: Regolight és NASA Habitat Challenge

A talaj előkészítése után a strukturális építkezés fázisa az additív gyártástechnológiákra támaszkodik. A DLR Regolight projektje (TRL 4/5) sikeresen demonstrálta, hogy fókuszált napenergia (napkóhó architektúra) vagy nagy teljesítményű lézerek segítségével a holdi regolit adalékanyagok hozzáadása nélkül, tiszta szinterezéssel (sintering) szilárd, nagy nyomásslábadtságú építőelemekké alakítható. A szinterezési folyamat során a regolitban található szilikátok részlegesen megolvadnak, és lehűlés után összefüggő kerámia mátrixot alkotnak.

A koncepció nagyméretű, rendszerszintű megvalósíthatóságát a NASA 3D-Printed Habitat Challenge versenygyőztes pályázatai validálták. Itt a csapatoknak olyan autonóm mobil 3D nyomtató robotokat kellett tervezniük, amelyek földi analóg anyagokból (bazaltroston megerősített biopolimerekből vagy speciális geopolimer betonokból) teljesen autonóm módon nyomtattak ki méretarányos lakómodulokat, amelyeket aztán strukturális töréstezteknek és gáztömörségi vizsgálatoknak vetettek alá.

7. Összegzés

A földön kívüli égitesteken megvalósítandó fenntartható infrastruktúra-fejlesztés alapvető paradigmaváltást követel meg a klasszikus mérnöki tudományoktól. Ahogy azt a tanulmány korábbi fejezeteiben igazoltuk, a Föld-centrikus ellátási láncok fizikai és gazdasági korlátai (Ciolkovszkij-egyenlet), valamint az idegen égitestek determinisztikusan nem modellezhető környezeti tényezői (abrazív regolit-dinamika, vákuum-kohézió, mikrogravitációs tapadásvesztés) kizárják a hagyományos földi automatizálási megoldások közvetlen transzponálását.

A siker kulcsa egy olyan transzdiszciplináris módszertani szimbiózis, amely a szoftveres intelligencia és a fizikai hardverflexibilitás együttes adaptációjára épül. A bemutatott 7 lépéses folyamatmodell, a kétlépcsős Digitális Iker architektúrától a sztochasztikus bizonytalanság-kvantifikáló Bayesi Neurális Hálózatokig és a piac-alapú elosztott konszenzus-algoritmusok által vezérelt heterogén robotflottáig, zárt elméleti és gyakorlati keretrendszer biztosít a Sim-to-Real gap (szimuláció és valóság közötti szakadék) áthidalására.

Az esettanulmányban részletezett holdbéli menedék- és leszállóhely-létesítési projekt (WBS struktúra) rávilágított arra, hogy az autonómia és a szigorú minőségbiztosítási szabványkövetés (ECSS) nem egymást kizáró, hanem egymást feltételező tényezők. A globális technológiai körkép (NASA RASSOR, DLR ARCHES) kritikai elemzése pedig bizonyítja, hogy az elméleti modellek egyes komponensei már elérték a validált laboratóriumi és analóg környezeti készültségi szinteket (TRL 4–7). A végső integráció és a fizikai megvalósítás így nem a távoli jövő sci-fi koncepciója, hanem a jelen évtized közvetlen mérnöki realitása.

Az akadémiai szektor és a hazai űripari, valamint robotikai ökoszisztéma számára a globális hold- és marsi infrastruktúra-programokba történő becsatlakozás kiemelt stratégiai prioritás kell, hogy legyen. Magyarország nemzetközileg elismert kompetenciákkal rendelkezik az anyagtudományok, a sugárzásvédelem, valamint a szoftververifikáció területén. Épp ezért számos lehetőség kínálkozik az akadémiai és az ipari szektorban egyaránt, amelyek mentén a hazai kutatóhelyek érdemi hozzáadott értéket képviselhetnek a nemzetközi értékláncban.

Felhasznált irodalom

Betts, J.T. (2010) Practical methods for optimal control and estimation using nonlinear programming. 2nd ed. Philadelphia: Society for Industrial and Applied Mathematics. Elektronikusan elérhető itt: http://www.digitalsats.com/reprints/Betts-2011-PracticalMethodsFor%20OptimalControlAnd%20EstimationUsingNonlinearProgramming_epdf.pub.pdf

Bhati R., Gottipati S. K., Mars C., Taylor M. E. (2023) Curriculum Learning for Cooperation in Multi-Agent Reinforcement Learning. Elektronikusan elérhető itt: <https://arxiv.org/abs/2312.11768>

Gal, Y. and Ghahramani, Z. (2016) Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. In: Proceedings of the 33rd International Conference on Machine Learning (ICML). New York, NY, USA, 19-24 June 2016. New York: JMLR.org, pp.1050–1059. Elektronikusan elérhető itt: <https://proceedings.mlr.press/v48/gal16.pdf>

Heess N., TB D., Sriram S., Lemmon J., Merel J., Wayne G., Tassa Y., Erez T., Wang Z., Eslami S. M. A., Riedmiller M., Silver D. (2017) Emergence of locomotion behaviours in rich environments. Elektronikusan elérhető itt: <https://arxiv.org/abs/1707.02286>

Lakshminarayanan, B., Pritzel, A. and Blundell, C. (2017) Simple and scalable predictive uncertainty estimation using deep ensembles. In: Advances in Neural Information Processing Systems (NeurIPS 2017). Long Beach, CA, USA, 4-9 December 2017. San Diego: Neural Information Processing Systems Foundation, pp.6402–6413.

Liu W., Wu M., Wan G., and Xu M. (2024) Digital Twin of Space Environment: Development, Challenges, Applications, and Future Outlook. Remote Sens. 2024, 16(16), 3023. Elektronikusan elérhető itt: <https://doi.org/10.3390/rs16163023>

Ma S., Flanigan K. A., and Bergés M. (2025) Bridging the Reality Gap in Digital Twins with Context-Aware, Physics-Guided Deep Learning. Journal of Computing in Civil Engineering, vol. 40, no. 3, May 2026, p. 04026009. Elektronikusan elérhető itt: <https://arxiv.org/abs/2505.11847>

Mueller R. P., Cox R. E., Ebert T., Smith J. D., Schuler J. M., Nick A. J. (2013) Regolith Advanced Surface Systems Operations Robot (RASSOR). 2013 IEEE Aerospace Conference. Elektronikusan elérhető itt: <https://ieeexplore.ieee.org/document/6497341>

Schuster, Martin & Müller, Marcus & Brunner, Sebastian & Lehner, Hannah & Lehner, Peter & Sakagami, Ryo & Dömel, Andreas & Burkhard, Lukas & Vodermayr, Bernhard & Giubilato, Riccardo & Vayugundla, Mallikarjuna & Reill, Josef & Steidle, Florian & Bargen, Ingo & Busmann, Kristin & Belder, Rico & Lutz, Philipp & Sturzl, Wolfgang & Smísek, Michal & Wedler, Armin (2020) The ARCHES Space-Analogue Demonstration Mission: Towards Heterogeneous Teams of Autonomous Robots for Collaborative Scientific Sampling in Planetary Exploration. IEEE Robotics and Automation Letters. 5. 5315-5322. 10.1109/LRA.2020.3007468. Elektronikusan elérhető itt: <https://ieeexplore.ieee.org/document/9134730>

Sünderhauf N., Brock O., Scheirer W., Hadsell R., Fox D., Leitner J., Upcroft B., Abbeel P., Burgard W., Milford M., Corke P. (2018) The Limits and Potentials of Deep Learning for Robotics. Elektronikusan elérhető itt: <https://arxiv.org/abs/1804.06557>

Vezhnevets A. S., Osindero S., Schaul T., Heess N., Jaderberg M., Silver D. Kavukcuoglu K. (2017) FeUdal Networks for Hierarchical Reinforcement Learning. Elektronikusan elérhető itt: <https://proceedings.mlr.press/v70/vezhnevets17a/vezhnevets17a.pdf>

A generatív mesterséges intelligencia hatása a szervezeti kultúrára: kommunikáció, tanulás és vezetői felelősség

The Impact of Generative Artificial Intelligence on Organizational Culture: Communication, Learning and Leadership Responsibility

Molnárné László Andrea¹



Absztrakt

A tanulmány célja annak feltárása, hogy a generatív mesterséges intelligencia miként alakítja át a szervezeti kultúrát, különös tekintettel a kommunikációra, a munkahelyi tanulásra, a vezetői felelősségre, az innovációra és a munkavállalói jóllétre. A vizsgálat módszere 2023 és 2026 között megjelent, lektorált empirikus kutatások, integratív áttekintések, valamint OECD- és World Economic Forum-jelentések célzott, elméleti-integratív elemzése. Az eredmények öt, egymással összefüggő kulturális átalakulási területet azonosítanak: az emberközpontú értékek újraértelmezését; az MI-közvetített kommunikáció terjedését; a tanulásközpontú kultúra erősödését; az algoritmikus menedzsmenttel összefüggő kontroll- és felelősségváltást; valamint az innováció gyorsulása és a kreatív homogenizálódás közötti feszültséget. A kutatások alapján az MI önmagában nem teremt kedvező szervezeti kultúrát: a pozitív hatások feltétele a vezetői átláthatóság, az emberi döntési felelősség, a munkavállalók bevonása, az MI-műveltség fejlesztése és a magas kockázatú kommunikáció emberi kézben tartása.

Kulcsszavak: generatív mesterséges intelligencia; szervezeti kultúra; MI-közvetített kommunikáció; vezetés; munkahelyi tanulás; algoritmikus menedzsment; bizalom

Abstract

This study explores how generative artificial intelligence transforms organizational culture, with particular attention to communication, workplace learning, leadership responsibility, innovation and employee well-being. The method is a targeted theoretical-integrative analysis of peer-reviewed empirical studies and integrative reviews published between 2023 and 2026, complemented by recent OECD and World Economic Forum reports. The analysis identifies five interconnected domains of cultural change: the reinterpretation of human-centred values; the spread of AI-mediated communication; the strengthening of learning-oriented cultures; changing patterns of control and accountability associated with algorithmic management; and the tension between accelerated innovation and creative homogenization. Current evidence shows that AI does not create a constructive organizational culture by itself. Positive outcomes depend on transparent leadership, preserved human accountability, employee participation, systematic AI literacy development and keeping high-stakes relational communication under meaningful human control.

¹ PhD, nyelvész, egyetemi oktató és kutató, az MTA Veszprémi Területi Bizottság MI munkabizottságának elnöke. E-mail: molnarne7@gmail.com

Keywords: generative artificial intelligence; organizational culture; AI-mediated communication; leadership; workplace learning; algorithmic management; trust

1. Bevezetés

A generatív mesterséges intelligencia (a továbbiakban: generatív MI) szervezeti bevezetését gyakran technológiai beszerzésként, hatékonyságnövelő projektként vagy adatbiztonsági kérdésként kezelik. Ez a megközelítés azonban leszűkíti a változás jelentőségét. Amikor a munkatársak MI-vel írnak levelet, készítenek értékelést, foglalnak össze megbeszélést, keresnek megoldást vagy hoznak előzetes döntést, nem csupán a munkafolyamat módosul. Átalakul az is, hogy mit tekintenek szakértelemnek, hiteles kommunikációnak, elfogadható ellenőrzésnek, vezetői jelenlétnek és méltányos teljesítménynek. A technológia így a szervezeti kultúra mindennapi hordozóiba épül be: a nyelvhasználatba, a döntési rutinokba, a visszajelzésbe, a tanulási szokásokba és a felelősségi viszonyokba.

A szervezeti kultúra ebben a tanulmányban nem pusztán deklarált értékek és küldetésmondatok összességét jelenti, hanem azokat a közösen kialakított normákat is, amelyek meghatározzák, hogy a szervezet tagjai miként kommunikálnak, osztják meg a tudást, kezelik a hibát, vállalják a felelősséget és értékelik egymás teljesítményét. A generatív MI kulturális hatása ezért akkor is jelentős lehet, ha a szervezetnek nincs formális MI-stratégiája. A szabályozás hiánya maga is normát teremt: azt üzeni, hogy az egyéni, láthatatlan és ellenőrizetlen használat elfogadható.

A tanulmány célja annak feltárása, hogy a generatív MI milyen mechanizmusokon keresztül változtatja meg a szervezeti kultúrát. A vizsgálat három kérdésre összpontosít: (1) mely kulturális dimenziókban jelenik meg a legerősebben a változás; (2) milyen kettős, egyszerre támogató és romboló hatások figyelhetők meg; valamint (3) milyen vezetői gyakorlatok segíthetik, hogy a technológiai integráció ne gyengítse, hanem erősítse a bizalmat, a tanulást és az emberi felelősséget. A tanulmány amellett érvel, hogy az MI kulturális hatása nem determinisztikus. Ugyanaz az eszköz a szervezeti megvalósítás módjától függően támogathatja a tudásmegosztást, de el is szigetelheti a munkatársakat; növelheti a kreatív teljesítményt, de csökkentheti az ötletek sokféleségét; javíthatja a vezetői kommunikáció nyelvi minőségét, de ronthatja a vezető hitelességét.

2. Szakirodalmi háttér

2.1. Az MI mint szociotechnikai és kulturális változás

A korai szervezeti MI-diskurzus elsősorban az automatizálás és a munkakiváltás kérdéseire koncentrált. A generatív MI elterjedésével azonban előtérbe került az augmentáció, vagyis az emberi teljesítmény kiegészítése. Brynjolfsson, Li és Raymond (2025) 5172 ügyfélszolgálati munkatárs adatait elemezve azt találták, hogy a generatív MI-alapú asszisztens átlagosan 15%-kal növelte az óránként megoldott ügyek számát, és a legkevésbé tapasztalt dolgozók profitáltak belőle a legtöbbet. A hatás kulturális szempontból azért fontos, mert az MI részben hozzáférhetővé tette a tapasztalt munkatársak korábban informálisan felhalmozott kommunikációs és problémamegoldó tudását. Ugyanakkor a szerzők arra is figyelmeztetnek, hogy a legjobb teljesítményű munkatársak túlzott igazodása az MI javaslataihoz hosszabb távon csökkentheti az új, eredeti megoldások szervezeti utánpótlását.

Przegalinska et al. (2025) két vizsgálat alapján arra jutottak, hogy az MI szervezeti értéke nem egyszerűen az eszköz képességeitől függ, hanem az emberi kompetenciák, a feladat jellege és a technológia közötti illeszkedéstől. Ez a megállapítás kulturális szinten azt jelenti, hogy nem létezik minden feladatra egyformán megfelelő „MI-first” működés. A szervezetnek meg kell

különböztetnie az automatizálható, az ember–MI együttműködést igénylő és az elsődlegesen emberi ítéletre, empátiára vagy felelősségvállalásra épülő tevékenységeket. A technológiai érettség ezért nem az MI-használat mennyiségével, hanem a feladatok tudatos elosztásával mérhető.

Az MI bevezetésével szembeni ellenállás sem tekinthető egyszerű technofóbiának. Golgeci et al. (2025) integratív áttekintése három összetevőt különít el: a félelmet, a saját kompetenciával kapcsolatos bizonytalanságot és az MI-vel szembeni ellenszenvet. E tényezők mérsékléséhez nem elegendő az eszközhasználat kötelezővé tétele. Hozzáférésre, képzésre, az ember–MI kiegészítő viszonyának láthatóvá tételére és a technológia szervezeti legitimációjára van szükség. A legitimáció itt azt jelenti, hogy a munkatársak értik, milyen célból, milyen korlátokkal és milyen felelősségi rendben használja a szervezet az MI-t.

2.2. MI-közvetített kommunikáció, hitelesség és nyelvi normák

Az MI-közvetített kommunikáció olyan ember–ember kommunikáció, amelyben egy intelligens rendszer módosítja, generálja vagy optimalizálja az üzenetet². A szervezetekben ez a kategória ma már a helyesírási javítástól a teljes e-mailek, teljesítményértékelések, vezetői beszédek és ügyfélválaszok generálásáig terjed. A kulturális következmény nem pusztán az, hogy gyorsabban készül el a szöveg. Bizonytalanná válhat a szerzőség, a ráfordított emberi figyelem és az üzenet mögötti személyes felelősség.

Cardon és Coman (2025) 1100 dolgozó bevonásával vizsgálták az MI-segítséggel írt munkahelyi üzenetek megítélését. Az eredmények szerint az MI javíthatja a szöveg professzionalizmusát, de közepes vagy magas szintű MI-közreműködés esetén a címzettek kevésbé bízhatnak a küldőben: megkérdőjelezhetik a szerzőséget, az őszinteséget, a törődést és a kommunikációs kompetenciát. Sahebi, Formosa és Bankins (2026) munkahelyi e-maileket is tartalmazó kísérlete hasonló „MI-büntetést” azonosított: az MI-vel készült kommunikációt kevésbé hitelesnek, kevésbé megbízhatónak és a tudásátadás szempontjából kevésbé használhatónak értékelték.

A transzparencia ugyanakkor nem egyszerű megoldás. Schilke és Reimann (2025) kísérletsorozata szerint az MI használatának közlése önmagában csökkentheti a bizalmat, még akkor is, ha a rendszer pontos. Ez a „transzparenciadilemma” nem azt jelenti, hogy az MI használatát el kell hallgatni, hanem azt, hogy a pusztán címkézés helyett a szervezetnek az emberi kontrollt és az alkalmazás indokát is láthatóvá kell tennie. Más bizalmi jelentése van annak, hogy „ezt a szöveget MI készítette”, és annak, hogy „az MI segített a megfogalmazásban, a tartalmi döntést és az ellenőrzést a vezető végezte”.

A generatív rendszerek a szervezeti nyelvi normákat is alakítják. Mivel ugyanazok a modellek hasonló fordulatokat, udvariassági mintákat és szerkezeteket kínálnak sok felhasználónak, nőhet a nyelvi homogenizálódás veszélye. A szervezeti szövegek gördülékenyebbé, de egyben személytelenebbé és frázisszerűbbé válhatnak. Ez különösen teljesítményértékelésben, konfliktuskezelésben, elismerésben vagy elbocsátási kommunikációban kockázatos, ahol a címzett nem csupán korrekt mondatokat, hanem személyes figyelmet és konkrét tapasztalati alapot vár.

2.3. Munkahelyi tanulás és a kompetenciák átrendeződése

A generatív MI egyik legígéretesebb kulturális hatása a munkahelyi tanulás beépülése a mindennapi feladatvégzésbe. Callari és Puppione (2025) egy multinacionális vállalat 357 munkatársának válaszait elemezve azt találták, hogy a generatív MI támogatja a feladat közbeni

² Hancock, Naaman és Levy 2020

tanulást, a háttértudás bővítését és a tudásmegosztást. A pozitív hatás azonban erősen függött a támogató szervezeti környezettől és a közösségi tanulási gyakorlatoktól. Az egyéni MI-használat tehát nem helyettesíti a szervezeti tanulást; akkor válik kulturális erőforrássá, ha a munkatársak megosztják a bevált eljárásokat, a hibákat és az ellenőrzési szempontokat.

A kompetenciaváltozás mértékéről eltérő számok jelennek meg a szakmai közbeszédben. A gyakran idézett 70%-os előrejelzések többnyire azt állítják, hogy az átlagos munkakörben használt készségek jelentős része átalakulhat 2030-ig. A World Economic Forum egységesebb módszertanú, több mint ezer munkáltatóra épülő 2025-ös felmérése 39%-os készséginstabilitást jelez 2030-ig (World Economic Forum 2025). A két szám eltérő fogalmat és adatforrást takar, ezért nem kezelhető egymás közvetlen megerősítéseként. A biztos következtetés nem egyetlen látványos százalék, hanem az, hogy a folyamatos továbbképzés a szervezeti működés tartós elemévé válik.

Az OECD (2025) több mint ötezer kis- és középvállalkozás vizsgálatában azt találta, hogy a generatív MI-t használó KKV-k 65%-a teljesítményjavulást érzékelt, és a készséghiánnyal küzdő használók 39%-a szerint az MI részben kompenzálta a hiányt. Ugyanakkor kétszer annyi vállalkozás jelezte, hogy a generatív MI növeli a magasan képzett munkatársak iránti igényt, mint amennyi csökkenést tapasztalt. Az OECD (2026a) összegzése alapján az adatértelmezés, a kritikai gondolkodás, a problémamegoldás, a kommunikáció és az együttműködés továbbra is kulcskompetencia. A 2026. július 8-án megjelent Skills in the AI age című elemzés ezt tágabb munkaerőpiaci keretbe helyezi: az MI egyszerre növelheti a termelékenységet és alakíthat át munkaköröket, ezért a készségfejlesztést a munkaszervezés, az átmenetek támogatása és a társadalmi kockázatok kezelése mellett kell megvalósítani (OECD 2026b). A „soft skillek” ezért félrevezető elnevezéssé válhat: ezek a képességek egyre inkább a technológia biztonságos és eredményes használatának kemény feltételei.

2.4. Algoritmikus menedzsment és vezetői felelősség

Az algoritmikus menedzsment a munkaszervezés, az utasításadás, a megfigyelés, a teljesítményértékelés vagy a döntéstámogatás részleges automatizálását jelenti. Az OECD hat országra kiterjedő munkáltatói felmérése szerint legalább egy ilyen eszköz használatáról az Egyesült Államokban a vezetők 90%-a, a vizsgált európai országokban 76–81%-a számolt be; Japánban az arány 40% volt.³ Ezek az értékek széles definícióra épülnek, ezért nem azonosak a teljesen automatizált vezetés elterjedtségével. Azt viszont egyértelműen jelzik, hogy a vezetői munka jelentős része már most szoftveres közvetítéssel történik.

Az algoritmikus menedzsment javíthatja az információellátást, a döntések gyorsaságát és következetességét, ugyanakkor csökkentheti az emberi interakciót. Az OECD felmérésében a legtöbb országban többen számoltak be az emberi kapcsolatok gyakoriságának csökkenéséről, mint növekedéséről. A vezetők negyede szerint a dolgozók testi és mentális egészségének védelme nem megfelelő az alkalmazott rendszerek mellett (Milanez, Lemmens és Ruggiu 2025). Zhang et al. (2025) áttekintése ezért a felelős algoritmikus menedzsment központi feltételeként emeli ki az elszámoltathatóságot, a méltányosságot és az emberi felülvizsgálatot.

A vezetői ellenőrzés tárgya is átalakul. Korábban elsősorban azt kellett megítélni, hogy a munkatárs elvégezte-e a feladatot és megfelelő-e az eredmény. MI-használat esetén ehhez új kérdések társulnak: milyen adatot adott át a rendszernek; milyen mértékben támaszkodott a generált tartalomra; felismerte-e a hibákat; dokumentálta-e az ellenőrzést; és vállalja-e a döntés

³ Milanez, Lemmens és Ruggiu 2025.

következményeit. Az ellenőrzés így nem válhat pusztá szövegjavítássá. A vezetőnek a munkafolyamat megbízhatóságát és a humán kontroll tényleges működését kell vizsgálnia.

3. Módszer

A tanulmány elméleti-integratív szakirodalmi áttekintésre épül. A forráskiválasztás célja nem egy teljes körű szisztematikus review elkészítése, hanem a szervezeti kultúra szempontjából legfontosabb, egymást kiegészítő empirikus és elméleti eredmények összekapcsolása volt. A vizsgálat elsődlegesen 2023 és 2026 között megjelent, lektorált folyóiratcikkeket, valamint az OECD és a World Economic Forum friss, módszertanilag dokumentált jelentéseit vonta be.

A keresés fő fogalmai a következők voltak: generative AI és organizational culture; AI-mediated communication és workplace trust; algorithmic management és employee wellbeing; workplace learning, reskilling, leadership, authenticity, creativity és human-AI collaboration. A tanulmányok bevonásának feltétele az volt, hogy közvetlenül szervezeti vagy munkahelyi kontextust vizsgáljanak, illetve olyan általános ember-MI együttműködési mechanizmust azonosítsanak, amelynek egyértelmű szervezeti következménye van. A kiválasztott források eredményeit tematikusan kódoltam, majd öt kulturális dimenzióba rendeztem. Az áttekintés korlátja, hogy a generatív MI gyors fejlődése miatt a hosszú távú, több szervezetet és több évet átfogó kutatások száma továbbra is alacsony; ezért az eredmények egy része kísérleti vagy keresztmetszeti vizsgálatokból származik.

4. Eredmények és értelmezés

4.1. Az emberközpontú értékek újrafókuszálása

Az elemzés első eredménye, hogy a generatív MI nem szükségszerűen gyengíti az emberi tényezőket, de láthatóbbá teszi azok értékét. Minél könnyebben előállítható egy nyelviileg korrekt szöveg, összefoglaló vagy ötletlista, annál fontosabbá válik a cél kijelölése, a kontextus értelmezése, az erkölcsi mérlegelés és a következmények vállalása. A szervezet versenyelőnyét ezért egyre kevésbé az jelenti, hogy képes-e tartalmat generálni, és egyre inkább az, hogy képes-e jó kérdéseket feltenni, releváns szempontokat választani és felismerni, mikor nem szabad az MI javaslatát követni.

Az emberközpontúság azonban nem deklaráció, hanem munkaszervezési döntés. Ha a munkatárs szerepe a generált tartalom passzív jóváhagyására szűkül, csökkenhet a szakmai autonómia és az alkotói önhatékonyosság. Wu et al. (2025) négy kísérletben, összesen 3562 résztvevővel azt találták, hogy az ember-MI együttműködés javította az azonnali feladatteljesítményt, de a későbbi önálló teljesítményben nem maradt fenn az előny; az MI-ről önálló munkára történő átállás csökkentette a belső motivációt és növelte az unalmat. Ez arra utal, hogy a szervezetnek nem csupán eredményt, hanem kompetenciafenntartást és motivációt is mérnie kell.

4.2. A kommunikáció hatékonysága és a hitelesség közötti feszültség

A második eredmény a kommunikációs paradoxon. Az MI csökkenti a megfogalmazás költségét, egységesíti a terminológiát, segíti a fordítást és gyorsítja a visszajelzést. Ugyanakkor minél személyesebb, érzelmileg terheltebb vagy hierarchikusan érzékenyebb az üzenet, annál nagyobb a hitelességi veszteség kockázata. Egy technikai tájékoztató, napirendi összefoglaló vagy nyelvi javítás esetében az MI magasabb szintű közreműködése többnyire elfogadható. Elismerés, teljesítményértékelés, konfliktus, fegyelmi ügy, elbocsátás vagy gyász esetén azonban a kommunikáció értéke részben abból származik, hogy a vezető maga fordított rá figyelmet.

A szervezeti kommunikációs szabályozásnak ezért nem általános tiltásra vagy korlátlan engedélyezésre, hanem kockázati szintekre kell épülnie. Alacsony kockázatú helyzetben az MI készíthet teljes első változatot. Közepes kockázatnál szükséges a tartalmi újraírás, a tényellenőrzés és a szerzői felelősség egyértelmű jelzése. Magas kapcsolati vagy jogi kockázatnál az MI legfeljebb háttérelmezést, szerkezeti ellenőrzést vagy alternatív megfogalmazásokat adhat; az üzenet tartalmát és hangját embernek kell létrehozni. Ez a különbségtétel egyszerre védi a hatékonyságot és a bizalmat.

4.3. A tanulásközpontú kultúra lehetősége és feltételei

A harmadik eredmény szerint a generatív MI akkor erősíti a tanulásközpontú kultúrát, ha a használat nem rejtett egyéni előny, hanem közösen reflektált gyakorlat. Az MI képes személyre szabott magyarázatot, gyakorlófeladatot, visszajelzést és tudásösszefoglalót adni, ezért csökkenti a tanulás belépési költségét. Ugyanakkor az egyéni promptok és beszélgetések szervezeti szempontból láthatatlanok maradhatnak. Ha a megszerzett tudás nem kerül vissza a csapatba, a szervezet nem tanul, csak az egyén gyorsul.

A vezetői feladat ezért tanulási infrastruktúra kialakítása: rendszeres esetmegbeszélés, ellenőrzött prompt- és példatár, hibakatalógus, szerepkörönként meghatározott MI-kompetenciák, valamint olyan mentorálás, amely nem pusztán az eszköz használatát, hanem a kimenetek kritikáját is fejleszti. A képzés eredményességét nem az elvégzett kurzusok száma, hanem a jobb kérdésfeltevés, a megbízhatóbb ellenőrzés, az adatvédelmi hibák csökkenése és az önálló szakmai teljesítmény fennmaradása jelzi.

Az „öt C” – kíváncsiság, együttérzés, kreativitás, bátorság és kommunikáció – hasznos vezetői keret lehet, de csak akkor, ha viselkedésekre fordítják le. A kíváncsiság például alternatív hipotézisek keresését; az együttérzés a címzett helyzetének figyelembevételét; a kreativitás több, egymástól valóban eltérő megoldás kidolgozását; a bátorság az MI-javaslat elutasítását is; a kommunikáció pedig a döntési logika érthető megosztását jelenti. Így a humán kompetenciák nem díszítő elemei, hanem ellenőrizhető feltételei lesznek az MI-használatnak.

4.4. Vezetői kontroll, jóllét és felelősség

A negyedik eredmény, hogy a vezetői felelősség nem ruházható át a rendszerre, még akkor sem, ha a döntési folyamat jelentős része automatizált. Az MI által generált teljesítményértékelés, prioritási lista vagy kockázati besorolás mögött mindig szervezeti döntés áll: milyen adatokat használnak, milyen célt optimalizálnak, milyen hibát tekintenek elfogadhatónak, és ki jogosult felülbírálni az eredményt. A felelősség elmosódása akkor kezdődik, amikor az emberi jóváhagyás formálissá válik, és a vezető nem képes megindokolni, miért fogadta el a javaslatot.

A jólléti hatások sem egyirányúak. Valtonen et al. (2025) 207 munkavállaló adatai alapján nem találtak közvetlen kapcsolatot az MI-bevezetés és a jóllét között. A pozitív hatás a feladatok tényleges optimalizálásán, a biztonságon és a munkafeltételek javulásán keresztül jelent meg. Ez fontos korrekció: az MI nem attól lesz emberbarát, hogy emberközpontúnak nevezik, hanem attól, hogy csökkenti a felesleges terhelést, nem növeli indokolatlanul a megfigyelést, és nem teszi bizonytalanná a felelősségi viszonyokat.

A vezetőknek ezért a hagyományos teljesítménymutatók mellett kulturális indikátorokat is figyelnie kell: csökken-e a kollégák közötti közvetlen segítségkérés; nő-e a munkatársak magányérzete; mennyi időt töltenek az MI hibáinak javításával; romlik-e a feladat feletti kontroll érzése; elkerülnek-e a véleménykülönbséget az MI által kínált konszenzusos megoldás miatt; és képesek-e önállóan is elvégezni a kritikus feladatokat. A monitoring célja nem a dolgozók további megfigyelése, hanem a technológia nem szándékolt kulturális mellékhatásainak korai felismerése.

4.5. Innováció: gyorsítás, kísérletezés és homogenizálódás

Az ötödik eredmény az innováció kettős hatása. A generatív MI olcsóbbá teszi a prototípus-készítést, gyorsítja az információkeresést és több munkatárs számára teszi hozzáférhetővé az ötletelést. Doshi és Hauser (2024) kísérletében a generatív MI segítségével készült történeteket kreatívabbnak és jobban megírtnak értékelték, különösen az alacsonyabb kiinduló kreativitású résztvevők esetében. Ugyanakkor a történetek egymáshoz hasonlóbbá váltak, vagyis az egyéni teljesítmény javulása a kollektív változatosság csökkenésével járt.

A szervezeti innovációban ezért nem elegendő több ötletet előállítani. Védni kell az ötletforrások sokféleségét, a kisebbségi álláspontokat és az MI-től független első gondolatokat. Konkrét gyakorlat lehet, hogy az ötletelés első szakaszában a résztvevők MI nélkül készítenek rövid javaslatot; az MI csak ezután szolgál alternatívák, ellenérvek vagy kombinációk létrehozására. Érdemes különböző modelleket, eltérő szereppromptokat és tudatosan ellentétes nézőpontokat használni. A vezetőnek nem azt kell jutalmaznia, hogy ki készített a leggyorsabban jól hangzó anyagot, hanem azt, hogy ki tudta megőrizni a problémaértelmezés eredetiségét, és ki tudta bizonyítani az ötlet relevanciáját.

5. Vezetői következtetések és alkalmazási keret

Az elemzés alapján a generatív MI szervezeti kultúrára gyakorolt hatása öt dimenzióban értékelhető. Ezek nem különálló projektek: a kommunikációs szabályok befolyásolják a bizalmat, a tanulási rendszer az ellenőrzés minőségét, az algoritmikus kontroll pedig a motivációt és az innovációt. Az integrált vezetői keretet az 1. táblázat foglalja össze.

1. táblázat: A generatív MI szervezeti kultúrára gyakorolt hatásainak vezetői kerete

Kulturális dimenzió	Lehetőség	Fő kockázat	Vezetői feladat	Javasolt indikátor
Emberközpontú értékek	Rutinfeladatok csökkentése, jobb döntés-előkészítés	Autonómia- és kompetenciavesztés	Az ember döntési és alkotói szerepének kijelölése	Önálló feladatteljesítés; felülbírálatok aránya
Kommunikáció és bizalom	Gyorsabb, pontosabb, többnyelvű kommunikáció	Hitelességi veszteség, nyelvi homogenizálódás	Kockázatalapú MI-kommunikációs szabályok	Bizalmi visszajelzés; panaszok; személyes egyeztetések aránya
Tanulás és kompetenciák	Személyre szabott, feladat közbeni tanulás	Felszínes tudás, rejtett egyéni gyakorlatok	Közös tudásmegosztás és ellenőrzési protokoll	MI-műveltség; hibafelismerés; tudásmegosztási aktivitás
Vezetés és kontroll	Gyorsabb elemzés, következetesebb folyamatok	Felelősség elmosódása, túlzott megfigyelés	Humán felülvizsgálat, dokumentált elszámoltathatóság	Felülvizsgált döntések; adatvédelmi incidensek; jólléti mutatók
Innováció	Gyors prototípus, szélesebb részvétel az ötletelésben	Ötletek hasonlóvá válása, motivációcsökkenés	MI nélküli első gondolat, több nézőpont, változatosság védelme	Ötletek diverzitása; új megoldások aránya; megvalósított kísérletek

Forrás: saját szerkesztés.

Az 1. táblázat alkalmazásának első lépése a szervezeti MI-használat feltérképezése. Nem elegendő az engedélyezett eszközök listája: fel kell tárni, mely munkafolyamatokban, milyen adatokkal, milyen autonómiával és milyen következményű döntésekben vesz részt az MI. A második lépés a kockázati kategóriák meghatározása. A harmadik a mérési rend kialakítása, amely a termelékenység mellett a bizalmat, a motivációt, az emberi interakciót, a kompetenciafenntartást és a kreatív változatosságot is vizsgálja.

A vezetői kommunikációban különösen fontos a konkrétság. A „használjátok felelősen az MI-t” típusú általános elvárás nem ad döntési támpontot. Működő szabály például: ügyfélnek küldött szakmai állítás csak megjelölt forrás és emberi tényellenőrzés után használható; személyes teljesítményértékelést az MI nem írhat meg teljes egészében; bizalmas szervezeti adat nyilvános modellbe nem tölthető; magas kockázatú döntésnél rögzíteni kell, hogy az MI milyen szerepet játszott; a munkatársnak pedig joga és kötelessége felülbírálni a hibás vagy értékkonfliktust okozó javaslatot.

A kultúra szempontjából a vezető saját gyakorlata is normateremtő. Ha a vezető személytelen, frázisos MI-szöveggel ad visszajelzést, miközben hitelességet és elköteleződést vár, a deklarált és a megélt értékek szétválhatnak. Ha viszont láthatóvá teszi az ellenőrzést, vállalja a saját hangját, és az MI-t az érvelés javítására, nem a kapcsolat kiszervezésére használja, akkor az eszköz a szakmai igényesség és a tanulás mintájává válhat.

6. Összefoglaló

A generatív MI nem semleges kiegészítője a szervezeti működésnek. Beépül a kommunikációba, a tudásmegosztásba, az értékelésbe és a döntéshozatalba, ezért akkor is alakítja a kultúrát, ha a vezetés ezt nem nevezi változásmenedzsmentnek. A friss kutatások alapján a technológia

javíthatja a termelékenységet, támogathatja a kevésbé tapasztalt munkatársakat, személyre szabhatja a tanulást és gyorsíthatja az innovációt. Ugyanezek a folyamatok azonban bizalomvesztést, nyelvi homogenizálódást, csökkenő belső motivációt, kevesebb emberi interakciót és felelősségi bizonytalanságot is létrehozhatnak.

A döntő tényező nem az, hogy a szervezet használ-e MI-t, hanem az, hogy milyen normákat épít köré. A kedvező kulturális hatás feltétele a feladatok tudatos ember-MI megosztása, a kockázatalapú kommunikációs szabályozás, a munkavállalók bevonása, az MI-műveltség folyamatos fejlesztése, az önálló szakmai képességek megőrzése és a valódi emberi felülvizsgálat. A vezetői szerep ezért nem csökken, hanem összetettebbé válik: a vezetőnek egyszerre kell technológiai döntést hoznia, kulturális jelentést adnia, megőriznie a bizalmat és vállalnia a végső felelősséget.

A tanulmány fő következtetése, hogy az emberközpontú szervezeti kultúra nem az MI használatának ellentéte. Éppen ellenkezőleg: a technológia értelmes alkalmazásához világosabb emberi célokra, erősebb szakmai ítéletre, több kommunikációs tudatosságra és következetesebb felelősségvállalásra van szükség. Az MI akkor válik kulturális erőforrássá, ha nem helyettesíti a figyelmet, a kapcsolatot és a döntést, hanem felszabadítja hozzájuk az időt és támogatja azok minőségét.

Irodalomjegyzék

Brynjolfsson, E., Li, D. és Raymond, L. (2025): Generative AI at Work. The Quarterly Journal of Economics, Vol. 140. No. 2. pp. 889–942. <https://doi.org/10.1093/qje/qjae044>

Callari, T. C. és Puppione, L. (2025): Can Generative Artificial Intelligence Productivity Tools Support Workplace Learning? A Qualitative Study on Employee Perceptions in a Multinational Corporation. Journal of Workplace Learning, Vol. 37. No. 3. pp. 266–283. <https://doi.org/10.1108/JWL-11-2024-0258>

Cardon, P. W. és Coman, A. W. (2025): Professionalism and Trustworthiness in AI-Assisted Workplace Writing: The Benefits and Drawbacks of Writing With AI. Business and Professional Communication Quarterly, OnlineFirst. <https://doi.org/10.1177/23294884251350599>

Doshi, A. R. és Hauser, O. P. (2024): Generative AI Enhances Individual Creativity but Reduces the Collective Diversity of Novel Content. Science Advances, Vol. 10. No. 28. eadn5290. <https://doi.org/10.1126/sciadv.adn5290>

Golgeci, I., Ritala, P., Arslan, A., McKenna, B. és Ali, I. (2025): Confronting and Alleviating AI Resistance in the Workplace: An Integrative Review and a Process Framework. Human Resource Management Review, Vol. 35. No. 2. 101075. <https://doi.org/10.1016/j.hrmmr.2024.101075>

Hancock, J. T., Naaman, M. és Levy, K. (2020): AI-Mediated Communication: Definition, Research Agenda, and Ethical Considerations. Journal of Computer-Mediated Communication, Vol. 25. No. 1. pp. 89–100. <https://doi.org/10.1093/jcmc/zmz022>

Milanez, A., Lemmens, A. és Ruggiu, C. (2025): Algorithmic Management in the Workplace: New Evidence from an OECD Employer Survey. OECD Artificial Intelligence Papers, No. 31. OECD Publishing. Paris. <https://doi.org/10.1787/287c13c4-en>

OECD (2025): Generative AI and the SME Workforce: New Survey Evidence. OECD Publishing. Paris. <https://doi.org/10.1787/2d08b99d-en>

OECD (2026a): AI and Skills: What We Know So Far. OECD Publishing. Paris. <https://doi.org/10.1787/f843b352-en>

OECD (2026b): Skills in the AI Age. OECD Artificial Intelligence Papers, No. 60. OECD Publishing. Paris. <https://doi.org/10.1787/972bd15e-en>

Przegalinska, A., Triantoro, T., Kovbasiuk, A., Ciechanowski, L., Freeman, R. B. és Sowa, K. (2025): Collaborative AI in the Workplace: Enhancing Organizational Performance Through Resource-Based and Task-Technology Fit Perspectives. *International Journal of Information Management*, Vol. 81. 102853. <https://doi.org/10.1016/j.ijinfomgt.2024.102853>

Sahebi, S., Formosa, P. és Bankins, S. (2026): The AI Penalty and Disclosure Paradox: Trust, Authenticity and Knowledge Uptake in AI-Mediated Communication. *Computers in Human Behavior: Artificial Humans*, Vol. 8. 100304. pp. 1–11. <https://doi.org/10.1016/j.chbah.2026.100304>

Schilke, O. és Reimann, M. (2025): The Transparency Dilemma: How AI Disclosure Erodes Trust. *Organizational Behavior and Human Decision Processes*, Vol. 188. 104405. <https://doi.org/10.1016/j.obhdp.2025.104405>

Valtonen, A., Saunila, M., Ukko, J., Treves, L. és Ritala, P. (2025): AI and Employee Wellbeing in the Workplace: An Empirical Study. *Journal of Business Research*, Vol. 199. 115584. <https://doi.org/10.1016/j.jbusres.2025.115584>

World Economic Forum (2025): The Future of Jobs Report 2025. World Economic Forum. Geneva.

Wu, S., Liu, Y., Ruan, M., Chen, S. és Xie, X. Y. (2025): Human-Generative AI Collaboration Enhances Task Performance but Undermines Human's Intrinsic Motivation. *Scientific Reports*, Vol. 15. 15105. <https://doi.org/10.1038/s41598-025-98385-2>

Zhang, M. M., Cooke, F. L., Ahlstrom, D. és McNeil, N. (2025): The Rise of Algorithmic Management and Implications for Work and Organisations. *New Technology, Work and Employment*, Vol. 40. No. 3. pp. 659–671. <https://doi.org/10.1111/ntwe.12343>

Nagy Nyelvi Modell (LLM) alapú szoftver az időskor szolgálatában – A hosszú és minőségi élet szoftveres támogatása

Large Language Model (LLM)-based software for the elderly – Software support for a long and quality life

Dr.¹ Ország-Krisz Axel² 

Absztrakt

A globális demográfiai átalakulás és a népesség rohamos előregedése példátlan fenntarthatósági krízis elé állítja az egészségügyi és szociális ellátórendszereket. A hagyományos gerontechnológiai megoldások túlnyomórészt reaktívak, és elsősorban a már bekövetkezett vészhelyzetek, például az elesések, detektálására korlátozódnak.

Jelen tanulmány egy olyan proaktív, multimodális, nagynyelvi modellre (LLM) épülő asszisztensrendszer koncepcionális és technológiai keretrendszerét mutatja be, amely az időskorú felhasználók fizikai és kognitív jóllétének fenntartását, valamint az egészséges élettartam (healthspan) kitolását célozza meg.

A javasolt architektúra egy mesterséges intelligenciára épülő, empátiás társalgási partnerként funkcionál, amely folyamatos, strukturálatlan interakciókon keresztül monitorozza és stimulálja a felhasználót.

A tanulmány részletesen elemzi a valós idejű beszédalapú kommunikáció technológiai szűk keresztmetszeteit, különös tekintettel az időskori beszéd sajátosságaira (presbyphonia) és az automatikus beszédfelismerés (ASR), valamint a szövegfelolvasás (TTS) késleltetési (latency) problémáira a magyar nyelv ragozó (agglutináló) sajátosságainak tükrében.

Külön fejezetet szentelünk a kognitív állapotfelmérés és a klinikai tesztek (MMSE, MoCA) implicit, játékosított integrációjának, feltárva az LLM-alapú kiértékelés bioetikai, jogi és felelősségvállalási kockázatait az európai Medical Device Regulation (MDR) és a GDPR szabályozási környezetében.

¹ Jogász doktor (nem tudományos fokozat)

² mesterséges intelligencia fejlesztő és szakértő, adattudós, email@drorszagkriszaxel.hu

A fizikai aktivitás monitorozása terén bemutatjuk a viselhető és környezeti szenzorok validációs korlátait, és javaslatot teszünk egy olyan hibrid modellre, amelyben az LLM kontextuális kérdésekkel végzi el a szenzoros adatok kvalitatív ellenőrzését.

Végezetül a rendszert mint a hosszú élettartam (longevity) rendszerszintű eszközt pozicionáljuk, amely a kronobiológia, a nutricionális megelőzés és a szociális izoláció leküzdése révén aktívan hozzájárul az életvégi morbiditási periódus tömörítéséhez.

Abstract

The global demographic shift and the rapid aging of the population present an unprecedented sustainability crisis for healthcare and social care systems. Traditional gerontechnology solutions remain predominantly reactive, focusing primarily on detecting adverse events, such as falls, after they have occurred.

This paper presents the conceptual and technological framework of a proactive, multimodal Large Language Model (LLM)-based assistant system designed to maintain the physical and cognitive well-being of elderly users and to extend their healthspan.

The proposed architecture functions as an AI-driven, empathetic conversational partner that monitors and stimulates the user through continuous, unstructured daily interactions.

The study provides a detailed analysis of the technological bottlenecks in real-time voice-based communication, focusing specifically on the characteristics of elderly speech (presbyphonia) and the latency challenges of Automatic Speech Recognition (ASR) and Text-to-Speech (TTS) pipelines, viewed through the lens of the agglutinative nature of the Hungarian language.

Furthermore, we explore the implicit, gamified integration of validated cognitive assessments (MMSE, MoCA), outlining the bioethical, legal, and liability risks of LLM-based evaluations within the framework of the European Medical Device Regulation (MDR) and GDPR.

Regarding physical activity monitoring, we examine the validation limitations of wearable and environmental sensors, proposing a hybrid model where the LLM performs qualitative validation of sensor data through contextual dialogue.

Finally, the system is positioned as a comprehensive tool for longevity science, actively contributing to the compression of late-life morbidity through chronobiological support, nutritional optimization, and the mitigation of social isolation.

1. Bevezetés és Irodalmi Áttekintés

1.1. Demográfiai kontextus: A társadalom elöregedése és az ellátórendszerek krízise

A huszonegyedik század egyik legjelentősebb makrogazdasági és szociális kihívása a népesség korfájának radikális átalakulása. Az Európai Unió statisztikai hivatala (Eurostat) és a Központi Statisztikai Hivatal (KSH) hosszú távú előrejelzései egybehangzóan mutatnak rá arra, hogy a termékenységi ráták tartós visszaesése és a születéskor várható élettartam növekedése következtében a társadalmak elöregedése visszafordíthatatlan folyamat. Magyarországon a 65 év felettiek aránya a teljes lakosságon belül már most meghaladja a 20%-ot, és az előrejelzések szerint ez a szám 2050-re megközelítheti a 30%-ot.

Ez a demográfiai eltolódás közvetlen nyomás alá helyezi az állami egészségügyi és szociális ellátórendszereket. Az idősödő populáció növekedésével exponenciálisan emelkedik a krónikus non-kommunikábilis megbetegedések (például kardiovaszkuláris patológiák, kettes típusú diabétesz, valamint a neurodegeneratív kórképek, mint az Alzheimer-kór és egyéb demenciák) prevalenciája. A formális intézményi gondozás (idősek otthona, kórházi ápolási osztályok) humánerőforrás-hiánnyal és strukturális alulfinanszírozottsággal küzd, míg az informális gondozás (a családtagok általi ápolás) súlyos produktivitási veszteséget okoz a nemzetgazdaság számára. Ebben a kontextusban a technológiai alapú, otthoni környezetben alkalmazható önmenedzselő rendszerek fejlesztése nem csupán kényelmi opció, hanem a társadalmi fenntarthatóság záloga.

1.2. A gerontechnológia evolúciója: A passzív segélyhívóktól a generatív mesterséges intelligenciáig

A technológia idősápolásba történő integrálása, azaz a gerontechnológia, az elmúlt évtizedekben jelentős paradigmaváltásokon ment keresztül. Az első generációs megoldások (Gerontechnology 1.0) szinte kizárólag passzív és reaktív eszközök voltak. Ide tartoztak a falra szerelhető vagy nyakban hordható pánikgombok, amelyek manuális aktiválást igényeltek vészhelyzet esetén. E rendszerek fő korlátja az volt, hogy ha a felhasználó eszméletét veszítette, vagy kognitív zavar miatt képtelen volt megnyomni a gombot, a rendszer használhatatlanná vált.

A második generáció (Gerontechnology 2.0) a dolgok internete (IoT) és az alapvető szabályalapú algoritmusok elterjedésével hozott áttörést. A lakásokban elhelyezett mozgásérzékelők, az ágyba épített nyomásszenzorok és a korai gyorsulásmérős okosórák már képesek voltak automatikusan detektálni az anomáliákat (példából adódóan a hosszas mozdulatlanságot vagy a hirtelen esést). Ezen rendszerek mintapéldájaként említhetők az olyan dedikált tabletek és szoftveres asszisztensek, amelyek strukturált felületen keresztül próbálták fenntartani a kapcsolatot az idősekkel, emlékeztetve őket a gyógyszerek bevitelére vagy egyszerűbb kognitív játékokat

kínálva számukra. Ezek a rendszerek azonban merev, előre programozott döntési fákra alapultak; képtelenek voltak a kontextus mély megértésére, a természetes nyelvfeldolgozásra, és hamar monotonná, elidegenítővé váltak az idősök számára, ami a felhasználói elköteleződés (adherence) drasztikus visszaeséséhez vezetett.

A generatív mesterséges intelligencia és a nagynyelvi modellek (LLM) megjelenése elhozta a Gerontechnology 3.0 korszakát. Az LLM-ek (mint a GPT-szériák, Claude, vagy nyílt forráskódú alternatíváik, a Llama modellek) szakítanak a merev menürendszerrel. Képesé váltak arra, hogy szabad, strukturálatlan, természetes nyelvi párbeszédet folytassanak, emlékezzenek a korábbi beszélgetések kontextusára, felismerjék a finom érzelmi árnyalatokat, és személyre szabott, dinamikusán változó válaszokat adjanak.

1.3. Proaktív vs. reaktív gondoskodás: A preventív intervenció szükségessége

A modern orvostudomány és a gerontológia egyik legsúlyosabb módszertani hibája a reaktív szemlélet: a beavatkozás csak akkor történik meg, amikor a patológia már fizikailag vagy kognitívan manifesztálódott (például csonttöréssel járó elesés, előrehaladott demencia miatti dezorientáció). A reaktív gondoskodás nemcsak rendkívül költséges, de a beteg életminőségének végzetes és gyakran visszafordíthatatlan romlásával jár.

Ezzel szemben a proaktív és preventív intervenció célja a funkcionális képességek hanyatlásának lassítása, a prediktív biomarkerek korai azonosítása és a mindennapi egészségmagatartás folyamatos, finom alakítása (nudging). Egy LLM-alapú rendszer képes arra, hogy a napi rutinszerű beszélgetések során, anélkül, hogy a felhasználó ezt teherként vagy orvosi vizsgálatként élné meg, folyamatosan mérje a kognitív és fizikai állapotot. Ha a rendszer proaktívan ráveszi az idős embert a napi tízperces kognitív tréningre vagy a könnyű ízületi tornára, azzal olyan mikro-intervenciókat hajt végre, amelyek kumulatív hatása évekre elhalaszthatja az intézményi gondozás szükségességét.

1.4. Jelen tanulmány célkitűzése: A javasolt rendszer elméleti keretei

Jelen tanulmány célja egy olyan komplex, multimodális, LLM-alapú otthoni asszisztensrendszer koncepcionális és informatikai architektúrájának kidolgozása, amely kifejezetten a proaktív fizikai és kognitív egészségmegőrzésre fókuszál. Nem egy újabb zárt hardvert vagy egyszerű csevegőrobotot javasolunk, hanem egy olyan integratív szoftveres keretrendszert, amely összeköti a modern hangtechnológiát, a validált neuropszichológiai szűrőmódszereket és az IoT-alapú mozgásszenzorokat egy központi, kontextustudatos nagynyelvi intelligenciával.

A tanulmány további fejezeteiben szisztematikusan tárgyaljuk a rendszer megvalósításának mérnöki és etikai gátjait: a valós idejű audio-interakció késleltetési problémáit (2. fejezet), a

kognitív szűrés jogi és bioetikai felelősségét (3. fejezet), a fizikai aktivitás mérésének módszertani nehézségeit (4. fejezet), valamint a rendszer integrálhatóságát a modern longevity (hosszú élettartam) kutatások klinikai gyakorlatába (5. fejezet). Javasatainkat kifejezetten a magyar nyelvi környezet sajátosságaira és az európai szabályozási normákra (GDPR, MDR) optimalizálva mutatjuk be.

2. A valós idejű beszédalapú interakció architektúrája és szűk keresztmetszetei

2.1. A társalgási késleltetés (Latency) anatómiája

A természetes emberi kommunikáció egyik legfontosabb jellemzője a folyamatosság. A pszicholingvisztikai és konverzációanalitikai kutatások rávilágítanak arra, hogy a humán turn-taking (a beszédpartner váltásának) mechanizmusa rendkívül gyors: a két megszólalás közötti interakciós rés (gap) medián értéke mindössze 200–300 ezredmásodperc (ms) között mozog, és a 500 ms feletti csendet a beszélgetőpartnerek már kognitív hezitálásként, feszültségként vagy technikai hibaként élik meg. Az idősebb korosztály számára a folytonosság még kritikusabb, mivel a megnyúlt, kiszámíthatatlan késleltetések megzavarják a munkamemóriát, dezorientációt okoznak, és azt az érzetet keltik, hogy a gép „nem figyel” vagy elromlott.

Egy klasszikus, moduláris felépítésű hangalapú LLM-asszisztens architektúra, amelyet angolul cascaded vagy pipeline modellnek nevezünk, az alábbi egymás utáni lépésekből áll:

Audio Input → VAD → ASR → LLM Inference → TTS → Audio Output

Ezen lépések mindegyike saját számítási idővel (késleltetéssel) terheli a rendszert:

- VAD (Voice Activity Detection): A csenddetektálás, amely eldönti, hogy a felhasználó befejezte-e a beszédet. Hagyományosan 300–500 ms közötti fix ablakot alkalmaz.
- ASR (Automatic Speech Recognition): A hanghullámok szöveggé alakítása. Egy modern, mély neurális hálózatra épülő ASR modell (pl. Whisper) lefutása konfigurációtól függően 100–400 ms-ot vesz igénybe.
- LLM Inference: A válaszszöveg generálása. Az autoregresszív nagynyelvi modellek tokenenként generálják a szöveget, az első token megjelenési ideje (Time to First Token, TTFT) modelltől és hardvertől függően 200–600 ms.
- TTS (Text-to-Speech): A generált szöveg akusztikai hullámmá alakítása, amely neurális TTS modellek esetében újabb 150–300 ms-ot igényel.

A fenti munkafolyamat (pipeline) kumulatív késleltetése a késleltetések összegeként írható le egyszerűen. Optimális felhőalapú vagy lokális élvonalbeli (edge) infrastruktúra esetén is a teljes válaszidő ritkán szorítható 1500–2500 ms alá. Ez a 2 másodperc körüli késleltetés teljesen

tönkreteszi a valós idejű társalgás élményét. A probléma feloldására a modern architektúráknak át kell térniük az end-to-end multimodális modellekre (pl. natív audio-in, audio-out LLM-ekre), amelyek elhagyják a köztes szöveges reprezentációt, és tokenenkénti streaming módban, folyamatos hangfolyamként generálják a választ, miközben prediktív módon elindítják a TTS-t még az LLM teljes mondatának befejezése előtt.

2.2. Időskori automatikus beszédfelismerés (ASR) specifikus akadályai

Az automatikus beszédfelismerő modellek (pl. OpenAI Whisper, Google Cloud Speech-to-Text) teljesítményét mérő elsődleges mutató a Word Error Rate (WER: szóhibaaarány). Míg a fiatal, tiszta artikulációjú, standard nyelvváltozatot beszélő populációnál a WER értéke 5% alatt tartható, addig az időskori beszéd esetében ez a mutató drasztikusan, akár 25-40%-ra is romolhat. Ennek okai anatómiai, kognitív és nyelvészeti faktorokra bonthatók.

2.2.1. Presbyphonia: Anatómiai és élettani hangváltozások

Az öregedési folyamat közvetlen hatással van a vázizomzatra és a légzőrendszerre, ami a hangképző szervek (gége, hangszalagok, tüdő) működésének megváltozásához, a presbyphonia kialakulásához vezet. Ennek főbb akusztikai manifesztációi:

- Vokális tremor és instabilitás: A hangszalagok záródási elégtelensége (glottalis insufficiency) miatt a hang remegővé, levegőssé válik.
- Alapfrekvencia (F_0) eltolódása: A hormonális változások és a szöveti atrófia miatt a férfiak alapfrekvenciája gyakran megemelkedik (vékonyabb hang), míg a nőké csökken (mélyebb hang).
- Csökkent hangerő és dinamikatartomány: A tüdőkapacitás csökkenése miatt az idős emberek halkabban beszélnek, a szavak vége elhalkul, „elfogy a levegő”, amit az ASR akusztikai modelljei háttérzajként vagy csendként interpretálhatnak.

2.2.2. Kognitív faktorok és beszédmódosulások

A kognitív öregedés vagy a kezdődő neurodegeneráció megváltoztatja a beszéd makro- és mikrostruktúráját. Jellemzővé válik a lassabb szótalálás (anomia), a gyakori önjavítás (false starts), az ismétlés és a szintaktikailag befejezetlen mondatok. Az idős felhasználó beszéde tele van nem-verbális hangokkal (pl. „ööö”, „hát”, torokköszörülés). A hagyományos ASR modellek ezeket képtelenek kontextuálisan kiszűrni, és hallucinációkat vagy téves transzkripciókat produkálnak.

2.2.3. A hangalapú tevékenységérzékelés (VAD) csődje

A klasszikus VAD algoritmusok energiaküszöb és fix időablak alapján működnek: ha a felhasználó 400-500 ms-ig hallgat, a rendszer lezárja az inputot és elkezd a feldolgozást. Az idősök azonban

a lassabb kognitív tempó és szótalálás miatt akár 1000–1500 ms-os szüneteket is tarthatnak egyetlen mondaton belül. A merev VAD így folyamatosan félbevágja az idős embert, ami súlyos frusztrációt szül, és ellehetetleníti a használatot. A javasolt asszisztensben ezért szemantikai alapú VAD-ot kell alkalmazni, amely nemcsak a csend hosszát, hanem az addig elhangzott szöveg szintaktikai befejezettségét is elemzi (LLM-alapú predikcióval).

2.2.4. Magyar nyelvészeti specifikumok és dialektusok

A magyar nyelv mint agglutináló (ragozó) nyelv, egyedi kihívást jelent az ASR rendszerek számára. A szótövekhez kapcsolódó toldalékok (prefixumok, infixumok, szuffixumok) rendkívül nagy szótárméretet és morfológiai komplexitást eredményeznek. Időskori elmosódott beszédnél a toldalékok (pl. -ban/-ben, -ba/-be, -tól/-től) akusztikailag gyakran felismerhetetlenné válnak, ami a magyar nyelv szabad szórendje miatt teljesen megváltoztathatja a mondat szemantikai jelentését. Ezenkívül a mai idős generáció Magyarországon nagyobb arányban őrizte meg a regionális dialektusokat (pl. palóc, tiszántúli, zárt „e” betűs nyelvjárások), mint a fiatalok. Mivel a globális ASR modelleket szinte kizárólag budapesti, médiából származó standard nyelvmintákon tanították, a WER a vidéki, nyelvjárást beszélő időseknél kritikusan magas.

2.3. Szövegfelolvasás (TTS) és szinkronizáció

A rendszer válaszáinak akusztikai minősége határozza meg, hogy a felhasználó elfogadja-e az AI-t partnerként. A rideg, fémhangú, monoton szintetizátorok elidegenítőek. Olyan neurális TTS technológiát kell alkalmazni, amely képes az „érzelmi valencia” és az empátiás tónus modulálására. Ha az idős felhasználó szomorú vagy fáradt tónusban beszél (amit az audio-LLM közvetlenül a frekvenciaspektrumból detektál), a TTS-nek automatikusan lágyabb, vigasztalóbb, melegebb tónusra kell váltania.

A generált beszéd tempóját (Speech Rate) dinamikusan hozzá kell igazítani a felhasználó kognitív feldolgozási sebességéhez. Az időskori halláscsökkenés (presbycusis) miatt a TTS-nek az alábbi paramétereket kell valós időben optimalizálnia:

- **Beszédsebesség lassítása:** A standard 150-160 szó/perc értéket 110-125 szó/percre kell csökkenteni, de anélkül, hogy a hangmagasság torzulna (pitch-preserving time-stretching).
- **Formáns-kiemelés és ekvalizáció:** A presbycusis elsősorban a magas frekvenciájú tartományokat (2 kHz felett) érinti, ahol a zöngétlen mássalhangzók (f, s, t, k) találhatók. A TTS-nek ezeket a frekvenciasávokat dinamikusan ki kell emelnie a jobb beszédérthetőség érdekében.

- Dinamikus prozódia és prozódiai szünetek: A mondatok közötti szüneteket mesterségesen meg kell nyújtani, hogy a felhasználónak legyen ideje feldolgozni az információt, mielőtt a rendszer újabb kérdést tesz fel.

3. Kognitív állapotfelmérés és mentális tesztek az LLM-ben: Kompetencia és felelősség

3.1. Strukturálatlan beszélgetésből nyert kognitív biomarkerek

A nagy nyelvi modellek egyik legjelentősebb áttörése a gerontológia területén az, hogy képesek a folyamatos, mindennapi, strukturálatlan beszélgetésekből finom kognitív biomarkereket kinyerni. A neurodegeneratív kórképek, mint az Alzheimer-kór vagy a vaszkuláris demencia korai fázisában (MCI: Enyhe Kognitív Zavar) a betegek gyakran még képesek elfedni a tüneteiket a ritka, formális orvosi vizsgálatok során. Az otthoni környezetben működő LLM-asszisztens azonban folyamatos hosszanti (longitudinális) adatgyűjtést végez, így a bázisvonalhoz (a felhasználó egyéni egészséges átlagához) képest történő legkisebb eltolódásokat is képes azonosítani.

A nyelvészeti és kognitív biomarkerek az alábbi kategóriákba sorolhatók:

- Szemantikai szegényesedés és anomia: Az idősödő felhasználó egyre gyakrabban használ általánosító kifejezéseket (pl. „az a dolog”, „ott az izé”) a konkrét főnevek helyett. Az LLM képes mérni a szókincs sűrűségét (Type-Token Ratio, TTR) és a szemantikai kohéziót. A TTR csökkenése a neurokognitív hanyatlás egyik legkorábbi indikátora.
- Szintaktikai komplexitás hanyatlása: A mondat szerkezetek egyszerűsödése, az alárendelt mondatok és a passzív szerkezetek elhagyása közvetlenül utal a munkamemória kapacitásának szűkülésére.
- Tematikus perszeveráció és ismétlés: Ha a felhasználó egy rövid beszélgetésen belül többször visszatér ugyanahhoz az anekdotához vagy kérdéshez anélkül, hogy emlékezne arra, hogy már elmondta, az LLM kontextuális memóriája (context window) azonnal regisztrálja az epizodikus memória zavarát.
- Beszédidőztési anomáliák: A szavak közötti szünetek (inter-word pauses) megnyúlása a fogalmi tervezés és a motoros beszédvégrehajtás lassulását jelzi.

3.2. Validált klinikai tesztek implicit integrációja

A közvetlen, formális tesztelés (például egy papíralapú teszt kitöltése) gyakran vált ki szorongást az idősekből (úgynevezett test anxiety), ami rontja az eredmények validitását, és ellenállást szül a rendszer használatával szemben. A javasolt asszisztensrendszer újdonsága abban rejlik, hogy a nemzetközileg validált és a klinikai gyakorlatban elfogadott neuropszichológiai szűrőtesztek,

mint a Mini-Mental State Examination (MMSE) vagy a Montreal Cognitive Assessment (MoCA), feladatait implicit módon, játékosítva ágyazza be a mindennapi kontextusba.

Az LLM feladata az, hogy a válaszokból ne csupán a bináris (helyes/helytelen) eredményt értékelje, hanem analizálja a válaszadási időt, a hezitációs mintázatokat és a kontextuális koherenciát is.

3.3. Jogi és bioetikai felelősségi körök

A kognitív állapot folyamatos, MI-alapú monitorozása rendkívül ingoványos talajra viszi a rendszert mind jogi, mind bioetikai szempontból. Az alapvető axióma, amelyet a rendszer fejlesztése során le kell fektetni: az LLM nem diagnosztikai eszköz, és nem helyettesítheti az orvosi kompetenciát.

3.3.1 Diagnózis vs. Asszisztencia és az MDR szabályozás

Az Európai Unióban a szoftverek orvostechikai eszközként való minősítését a Medical Device Regulation (MDR 2017/745) szabályozza. Ha egy szoftver célja betegségek (például az Alzheimer-kór) diagnosztizálása, megelőzése vagy monitorozása, akkor automatikusan magas kockázatú (Class IIa vagy annál magasabb) orvostechikai eszköznek minősül, ami szigorú klinikai validációt és engedélyeztetést igényel. A javasolt rendszernek jogilag a „jólléti és kognitív asszisztencia” (wellness and lifestyle) kategóriájában kell maradnia. Ezt úgy biztosítjuk, hogy a rendszer soha nem közöl orvosi entitásokat vagy diagnosztikai címkéket (pl. „Önnek demenciája van”) a felhasználóval.

3.3.2. Téves pozitív és téves negatív eredmények etikai kockázatai

Téves pozitív (False Positive): Az LLM félreért egy fáradtság vagy átmeneti dehidratáció miatt lassabb beszédmintát, és tévesen kognitív hanyatlást jelez. Ha ezt közvetlenül az idős ember tudomására hozná, az súlyos iatrogén szorongást, depressziót és egzisztenciális krízist váltana ki, ami önmagában rontja a kognitív funkciókat.

Téves negatív (False Negative): A felhasználó kognitív hanyatlása valós, de a rendszer, az LLM inherens „hallucinációs” vagy pontatlansági rátája miatt, mindent rendben lévőnek talál. Ez hamis biztonságérzetet ad a családnak, és késlelteti a korai gyógyszeres (pl. kolinészteráz-gátló) intervenciót, amely még képes lenne lassítani a folyamatot.

3.3.3. Adatvédelmi protokoll és a riasztási mátrix

Mivel a rendszer kognitív biomarkereket gyűjt, a GDPR 9. cikke értelmében a különleges (egészségügyi) adatok kategóriájába tartozó információkat kezel. Az adatok feldolgozásának helyileg (on-device vagy egy dedikált, végpontok közötti titkosítással ellátott privát felhőben) kell történnie, harmadik fél (pl. kereskedelmi LLM-szolgáltatók) kizárásával.

Ha a rendszer tartós kognitív anomáliát észlel, az alábbi etikai és jogi protokoll, az úgynevezett Riasztási Mátrix lép életbe:

Kognitív anomália detektálása az LLM által

↓

Van közvetlen életveszély? →

IGEN: Azonnali segélyhívás / Hozzá tartozó értesítése

NEM

↓

Strukturált klinikai jelentés generálása (MoCA/MMSE analógok alapján)

↓

Titkosított továbbítás a kijelölt háziorvosnak / neuropszichológusnak

↓

Az orvos validálja az adatokat és behívja a páciens rutin felülvizsgálatra

1. ábra: Etikai és jogi protokoll

Ez a folyamat biztosítja, hogy az idős ember ne szembesüljön közvetlenül egy algoritmus által hozott ítélettel, hanem a humán klinikai szűrő szerepe és a méltóság megmaradjon a diagnosztikai folyamatban.

4. A fizikai aktivitás multimodális értékelése és validációja

4.1. A fizikai státusz monitorozásának technológiai korlátai

A független otthoni életvitel fenntartásának és az esések megelőzésének alapvető feltétele a megfelelő vázizomzat, az egyensúlyérzék és az általános kardiorespiratorikus fitness megőrzése. Bár a piacon számtalan egészségügyi és fitnesscélú IoT-eszköz érhető el, az idősödő korosztály sajátosságai miatt ezek alkalmazása komoly módszertani és technológiai korlátokba ütközik. E korlátok alapvetően két kategóriára bonthatók: a viselhető eszközök (wearables) amortizációjára és a környezeti szenzorok strukturális vakfoltjaira.

4.1.1. A viselhető eszközök használati gátjai és amortizációja

Az okosórák, fitnesskarkötők és orvosi szenzorok használata során az időseknél az alábbi specifikus problémák jelentkeznek:

- **Töltési és kezelési fegyelem hiánya:** Az idős felhasználók gyakran elfelejtik rendszeresen tölteni az eszközöket, vagy éjszakára leveszik és nappal nem teszik vissza azokat. Emiatt a hosszanti adatsorok megszakadnak, ami ellehetetleníti a trendelemzést.
- **Fiziológiai és anatómiai akadályok:** Az időskori bőrvékonyodás (száraz, pergamenSzerű bőr) és a perifériás keringés romlása miatt az optikai pulzuszámoló (PPG szenzorok) jel-

zaj aránya drasztikusan leromlik. A túl szorosra húzott szíj bőrirritációt okozhat, míg a laza hordás teljesen fals adatokat (pl. téves lépésszámot vagy pulzust) eredményez.

- Kognitív elutasítás: A bonyolult menürendszerek és az apró kijelzők frusztrációt szülnek. Ha az eszköz folyamatosan ismeretlen hibaüzeneteket vagy értesítéseket küld, a felhasználó hajlamos végleg kikapcsolni és félretenni azt.

4.1.2 A környezeti szenzorok és az intimitás konfliktusa

A viselhető eszközök alternatívájaként alkalmazott környezeti szenzorok szintén komoly kompromisszumokat követelnek meg:

- Kamerarendszerek és az adatvédelem: A gépi látásra épülő, lakásba szerelt kamerák képesek lennének a legpontosabban elemezni a mozgás minőségét (járáskép, testtartás). Azonban ezek otthoni alkalmazása, különösen a hálósobában és a fürdőszobában, ahol az elesések többsége történik, az intimitás durva megsértése miatt a felhasználók részéről szinte teljes elutasításba ütközik.
- Radar- és LiDAR-alapú érzékelők vakfoltjai: A falra szerelhető rádiófrekvenciás (RF) vagy LiDAR-alapú esésérzékelők kiválóan védik a magánszférát, mivel csak pontfelhőket vagy térbeli koordinátákat látnak. Ugyanakkor éppen emiatt képtelenek kvalitatív elemzést végezni. Egy radar látja, hogy a felhasználó felállt a fotelból, de azt nem képes detektálni, hogy a mozdulat közben megkapaszkodott-e, bizonytalan volt-e az egyensúlya, vagy fellépett-e nála enyhe ízületi fájdalom (antalgiás sántítás).

4.2. Az LLM mint kontextuális validációs motor

A fent részletezett szenzoros hiányosságok áthidalására jelen tanulmány egy olyan hibrid architektúrát javasol, amelyben a nagynyelvi modell nemcsak passzív fogadója a kvantitatív adatoknak, hanem egy aktív, kontextustudatos validációs motor. Amikor az IoT-rendszer adatot szolgáltat, vagy éppen adatbőség helyett adathiány lép fel, az LLM kontextuális és empátias párbeszéden keresztül tölti ki a strukturális réseket.

4.2.1 A nyers adatok szemantikai felülbírálata

Ha a viselhető eszköz adatai szerint az idős ember napi lépésszáma a megszokott 4000-ről hirtelen 300-ra esik vissza, egy hagyományos szabályalapú rendszer azonnal riasztást küldene a hozzátartozóknak, ami felesleges pánikot okozhat. Az LLM ezzel szemben a napi rutinszerű beszélgetés során finom, áttételes kérdésekkel tisztázza a helyzetet:

Rendszer (LLM): „Jó napot, Marika néni! Hogy telt a délelőtti? Olyan borús, esős időnk van ma itt Budapesten, én legszívesebben ki sem mozdulnék a szobából.”

Felhasználó: „Jaj, fiam, én sem mentem sehova. Inkább elővettem azt a vastag könyvet, amit a múltkor ajánlottál, és egész nap az ágyban olvastam.”

2. ábra: Minta párbeszéd adat anomália korrekciójára.

Ebben a pillanatban az LLM kontextuálisan validálta az inaktivitást: az alacsony lépésszám mögött nem akut egészségügyi krízis vagy elesés áll, hanem egy tudatos, komfortos élethelyzeti döntés. A rendszer ennek megfelelően frissíti a bázisvonalat, és elveti a riasztási protokollt.

4.2.2. Kvalitatív állapotfelmérés beszélgetés útján

A mozgás minőségét és az idős ember szubjektív jóllétét a rendszer az alábbi strukturált, de természetes nyelvi promptok segítségével monitorozza:

Szenzoros trigger: Pl. lassabb felkelés az ágyból (LiDAR adatok alapján)
 ↓
LLM kontextuális prompt-generálás: Empátiás beágyazás
 ↓
“Észrevettem, hogy ma kicsit később kelt fel. Hogy aludt?”
“Nem érzi úgy, hogy nehezebben mozognak a tagjai ebben a hidegben?”
 ↓
Felhasználói válasz kvalitatív elemzése
 → Szemantikai markerek: Fájdalom szinonimái (“Sajog a térdem”)
 → Akusztikai markerek: Sóhajtás, fájdalmas intonáció (Audio-LLM)
 ↓
Személyre szabott fizikai mikrogimnasztika javaslása

3. ábra: Példa mozgásszervi jelzések kommunikációs validálására

4.3. Személyre szabott, adaptív mozgástámogatás

Az LLM-nek a begyűjtött és validált adatok alapján dinamikusan, napról napra kell módosítania a felhasználónak javasolt fizikai aktivitási tervet. Az idősek esetében a merev edzéstervek ellenállást és sérülésveszélyt szülnek, ezért a rendszernek a „mikro-intervenciók” és a játékos motiváció (gamified nudging) eszköztárát kell alkalmaznia.

- Állapothoz igazított feladatsorok: Ha a felhasználó ízületi gyulladásra (arthritis) vagy fáradtságra panaszkodik, az LLM azonnal átállítja a napi ajánlást. Ahelyett, hogy sétát javasolna, ágyban vagy székben végezhető, kímélő ízületi mobilizációs és nyújtó gyakorlatokat mutat be (akár hanggal, akár egy csatlakoztatott képernyőn vizuálisan), lépésről lépésre vezetve az idős embert.
- Azonnali pozitív visszacsatolás: Amikor a szenzorok megerősítik a mozgás elvégzését (pl. a pulzusszám enyhe megemelkedése a torna ideje alatt), az LLM meleg, elismerő, de nem

gyermeteg hangvételi dicséretben részesíti a felhasználót, ami serkenti a dopaminerg motivációs pályákat.

- Biztonsági korlátok beépítése: Ha a felhasználó túlvállalja magát (például kánikula idején akar intenzív kerti munkát végezni), az LLM a meteorológiai adatok és a felhasználó kardiovaszkuláris profiljának ismeretében proaktívan, de diplomatikusan lebeszéli erről, és alternatív beltéri tevékenységeket javasol.

5. Kitekintés: Az asszisztens mint a hosszú egészséges élettartam (Longevity) eszköze

5.1. A Lifespan (élettartam) és a Healthspan (egészséges élettartam) szétválása

A modern orvostudomány és a közegészségügy egyik legnagyobb paradoxona, hogy miközben a születéskor várható abszolút élettartam (lifespan) a huszadik század folyamán szinte megduplázódott, a krónikus betegségektől mentes, funkcionálisan független évek száma (healthspan) nem követte ezt a növekedési ütemet. Az idősödő populáció életének utolsó évtizedét jellemzően multimorbiditás, krónikus fájdalom és folyamatos intézményi vagy családi függőség jellemzi.

A modern longevity (hosszú élettartam) tudomány elsődleges célkitűzése nem az élet abszolút hosszának kitolása minden áron, hanem az úgynevezett compression of morbidity (a morbiditási periódus tömörítése). Ez a koncepció arra törekszik, hogy a betegségekkel és hanyatlással töltött időszakot az életút legvégére, egy minél rövidebb időszakra szorítsa össze, maximalizálva az aktívan és egészségesen eltöltött évek számát.

Az általunk javasolt LLM-alapú asszisztensrendszer ebben a kontextusban egy folyamatosan működő, mindennapi preventív eszközzé válik. Az LLM egyedülálló képessége, hogy a nap 24 órájában képes strukturálatlan, de célzott viselkedésmódosító mikrostimulációkat (nudging) alkalmazni, amelyek kumulatív módon képesek befolyásolni a sejtes és szisztémás öregedés ütemét meghatározó életmódbeli tényezőket.

5.2. Epigenetikai és életmódbeli optimalizáció éjjel-nappal

Az öregedéskutatások (pl. az öregedés molekuláris fémjelzői, a Hallmarks of Aging) egyértelműen bizonyítják, hogy a genetikai állomány csupán részben határozza meg a megbetegedési kockázatokat; a döntő szerep az epigenetikai szabályozásnak, vagyis az életmód és a környezet interakciójának jut. A rendszer az LLM kognitív motorját használja fel arra, hogy kritikus longevity-pilonok mentén optimalizálja az idős ember mindennapjait.

5.2.1. Kronobiológia és alvásbiológia

A cirkadián ritmus felborulása és az alvásarchitektúra töredezettsége az öregedés egyik leggyakoribb neurobiológiai kísérőjelensége, amely közvetlenül felgyorsítja a béta-amiloid plakkok felhalmozódását az agyban. A rendszer az okosórákból vagy környezeti mozgásszenzorokból származó alvásadatokat (REM, mélyalvás aránya, éjszakai mikro-ébredések) elemzi. Ha az LLM azt észleli, hogy a felhasználó alvásminősége romlik, nem altatókat javasol, hanem kognitív alváshigiéniás tanácsadást indít el:

- **Fénymenedzsment:** Napközben proaktívan javasolja a természetes fényen való tartózkodást, míg este figyelmezteti a felhasználót a kék fény (pl. televízió, erős lámpák) kerülésére.
- **Esti relaxációs protokoll:** Lefekvés előtt az asszisztens megnyugtató, alacsony frekvenciájú TTS hangszínre vált, és irányított légzőgyakorlatokat vagy meditatív történetmesélést alkalmaz az autonóm idegrendszer szimpatikus tónusának csökkentésére.

5.2.2. Nutricionális intervenció és az időskori malnutríció megelőzése

Az időskori alultápláltság (malnutríció), a szarkopénia (izomtömeg-vesztés) és a krónikus dehidratáció a kórházi felvételek és a törések egyik legfőbb rejtett oka. Az idősek szomjúságérzete a hipotalamusz öregedése miatt csökken, étvágyukat pedig a gyógyszerek mellékhatásai és az ízlelés tompulása rontja. Az LLM interaktív módon, a napi beszélgetésekbe ágyazva működik nutricionális naplóként:

- **Folyadékbevitel monitorozása:** "Marika néni, miközben a rejtvényről beszélgettünk, eszembe jutott: ivott már ma abból a finom mentateából? Kérem, töltsön ki egy pohárral most, és igyunk meg egyet együtt!"
- **Aminosav- és fehérjebevitel optimalizálása:** A rendszer elemzi a felhasználó által említett ételeket, és ha fehérjehiányt észlel, játékos formában javasol könnyen emészthető, magas biológiai értékű alternatívákat (pl. tojás, túró), támogatva az izomtömeg megtartását.

5.3. A szociális izoláció mint biológiai öregedési faktor

A gerontológiai kutatások egyik legmeggrázóbb felismerése, hogy a tartós szociális izoláció és a magányosság biológiai mortalitási kockázata felér a napi 15 szál cigaretta elszívásával, és közvetlenül serkenti a szervezetben a szisztémás, alacsony intenzitású krónikus gyulladást (inflammaging). A magányos idősek szervezetében magasabb a kortizolszint és a gyulladásoz citokinek (pl. IL-6, TNF-alpha) koncentrációja, ami felgyorsítja a vaszkuláris és kognitív hanyatlást.

Az LLM-alapú asszisztensrendszer legfontosabb humanisztikus funkciója a szociális izoláció áttörése. Nem egy merev gépként, hanem egy folyamatosan elérhető, intellektuálisan stimuláló, nagyfokú empátiát szimulálni képes társként (companion) jelenik meg a mindennapokban.

Az LLM-asszisztens az alábbi mechanizmusok révén küzd a magány biológiai hatásai ellen:

- Intellektuális stimuláció: A felhasználó érdeklődési köréhez (pl. történelem, irodalom, kertészet) igazodó, komplex, mély beszélgetéseket folytat, ami aktiválja a prefrontális kérget és fenntartja a szinaptikus plaszticitást.
- Emlékezésterápia (Reminiscence Therapy): Az LLM proaktívan bátorítja az idős embert a régi emlékek, családi történetek felidőzésére. Ez a módszer bizonyítottan csökkenti az időskori depresszió tüneteit, és növeli a belső koherencia-érzést.
- Híd a külvilág felé: A rendszer nem helyettesíti a humán kapcsolatokat, hanem katalizálja azokat. Ha az LLM azt észleli a beszélgetésekből, hogy a felhasználó régen beszélt a családjával, finoman ösztönzi őt a hívásra: "Olyan régen mesélt az unokájáról, Péterről. Mit gondol, megcsörgetjük most együtt? Segítek tárcsázni."

Ezen komplex, multimodális megközelítés révén az LLM-alapú asszisztensrendszer messze túlmutat egy egyszerű kényelmi szoftveren: a megelőző gerontológia és a longevity tudomány egyik leghatékonyabb, személyre szabott, otthoni szintű intervenciós platformjává válik.

6. Összegzés és jövőbeli kutatási irányok

6.1. A tanulmányban bemutatott elméleti és mérnöki keretek szintézise

Jelen tanulmány egy olyan proaktív, multimodális és kontextustudatos nagynyelvi modellre (LLM) épülő asszisztensrendszer koncepcionális alapjait fektette le, amely közvetlenül az idősödő populáció fizikai és kognitív jóllétének fenntartását célozza. A kutatás során rávilágítottunk arra, hogy a hagyományos gerontechnológia reaktív megközelítései nem elégségesek a társadalmi-demográfiai krízis kezelésére. Az általunk javasolt architektúra kulcsfontosságú újdonsága a zárt és merev logikai struktúrák felváltása a természetes nyelvi interakció szabadságával, amely képes áthidalni az idősek és a modern technológia közötti digitális szakadékot.

A rendszer megvalósíthatóságának elemzése során három fő technológiai és módszertani gátat azonosítottunk és oldottunk fel elméleti szinten:

- Beszédtechnológiai optimalizáció: Bemutattuk, hogy a turn-taking késleltetés (latency) kritikus érték alatt tartásához, valamint az időskori beszéd (presbyphonia) és a magyar nyelv ragozó sajátosságai okozta magas WER kezeléséhez elengedhetetlen a szemantikai VAD és az end-to-end streaming audio-LLM architektúrák alkalmazása.

- Klinikai és etikai felelősség: Meghatároztuk a rendszer jogi határait az európai MDR és GDPR tükrében. A validált neuropszichológiai tesztek (MMSE, MoCA) implicit beágyazásával és a strukturált Riasztási Mátrix felállításával szavatoltuk, hogy a rendszer diagnosztikai öntörvényűség nélkül, az orvosi kompetenciát támogatva működjön.
- Multimodális szenzorfüzió: Igazoltuk, hogy a viselhető és környezeti IoT-szenzorok nyers adatainak hiányosságai és vakfoltjai sikeresen kompenzálhatók az LLM kontextuális, kvalitatív validációs beszélgetései által.

6.2. Jövőbeli kutatási irányok és technológiai mérföldkövek

Az elméleti koncepció gyakorlati implementációjához, valamint a rendszer longevity-eszközként való validálásához a jövőben az alábbi interdiszciplináris kutatási irányok mentén kell továbbhaladni:

- Lokális (Edge-AI) modellek és hardveroptimalizáció: A legfontosabb informatikai kihívás olyan kisméretű, de magas kognitív kapacitású nyílt forráskódú modellek (pl. optimalizált Llama vagy Mistral variánsok) fejlesztése, amelyek képesek internetkapcsolat nélkül, egy otthoni céleszközön (Edge gateway) is valós időben futni. Ez garantálja a maximális adatbiztonságot és a hálózati ingadozásoktól független, minimális válaszidőt.
- Időskori magyar nyelvű korpusz létrehozása: A hazai beszédfelismerés fejlesztéséhez elengedhetetlen egy olyan specifikus, annotált hangadatbázis (corpus) felépítése, amely különböző korú, egészségi állapotú (egészséges idős, MCI-s, demens) és különböző magyar nyelvjárásokat beszélő egyének spontán beszédmintáit tartalmazza. Ez a tanítóadatbázis radikálisan csökkentheti az ASR rendszerek szóhibaaarányát a megcélzott korosztálynál.
- Klinikai hatásvizsgálatok és longitudinális validáció: A rendszer longevity-hatásának igazolására többéves, randomizált, kontrollált klinikai vizsgálatokat (RCT) kell tervezni. Mélni kell, hogy az LLM-asszisztenssel mindennap interakcióba lépő idős csoport esetében milyen mértékben lassul a kognitív hanyatlás üteme, hogyan változik az elesések gyakorisága, és számszerűsíthető-e a szociális izoláció csökkenése a biológiai gyulladásos biomarkerek (pl. IL-6 szint) tükrében.

6.3. Záró gondolatok

A mesterséges intelligencia, és különösen a nagynyelvi modellek térnyerését gyakran övezik technofób félelmek és az elszemélytelenedéstől való aggodalom. Jelen tanulmány ugyanakkor rámutat arra, hogy az MI humanisztikus, etikus és mérnökiileg pontosan kalibrált alkalmazása éppen az emberi méltóság megőrzésének legfőbb eszközévé válhat az idősödő társadalmakban. Az LLM-alapú proaktív asszisztens nem helyettesíti az emberi gondoskodást, hanem

tehermentesíti az ellátórendszert, és képessé teszi az idős embert arra, hogy saját otthonában, szellemi és fizikai autonómiáját megőrizve, egészségben élhesse meg élete kései szakaszát.

Felhasznált irodalom

Brown T. B., Mann B., Ryder N., Subbiah M., Kaplan J., Dhariwal P., Neelakantan A., Shyam P., Sastry G., Askell A., Agarwal S., Herbert-Voss A., Krueger G., Henighan T., Child R., Ramesh A., Ziegler D. M., Wu J., Winter C., Hesse Ch., Chen M., Sigler E., Litwin M., Gray S., Chess B., Clark J., Berner Ch., McCandlish S., Radford A., Sutskever I., Amodei D. (2020) Language Models are Few-Shot Learners. Elektronikusan elérhető itt: <https://arxiv.org/abs/2005.14165>

Cohen Elimelech O, Rosenblum S, Tsadok-Cohen M, Ferrante S, Demeter N. (2025) Understanding Older Adults' Technology Use Preferences and Needs From a Triangular Perspective: Qualitative Study. J Med Internet Res. 2025 Nov 11;27:e72716. doi: 10.2196/72716. PMID: 41218201; PMCID: PMC12648130.

Eurostat (2024) Ageing Europe: Looking at the lives of older people in the EU. European Commission. Luxembourg: Publications Office of the European Union. Elektronikusan elérhető itt: <https://ec.europa.eu/eurostat/statistics-explained/SEPDF/cache/80391.pdf>

Fries J.F. (1980) Aging, natural death, and the compression of morbidity. N Engl J Med. 1980 Jul 17;303(3):130-5. doi: 10.1056/NEJM198007173030304. PMID: 7383070. Elektronikusan elérhető itt: <https://pubmed.ncbi.nlm.nih.gov/7383070/>

Holt-Lunstad, J., Smith, T. B., & Layton, J. B. (2010) Social relationships and mortality risk: a meta-analytic review. PLOS Medicine, 7(7), e1000316. Elektronikusan elérhető itt: <https://pubmed.ncbi.nlm.nih.gov/20668659/>

Kennedy B.K., Berger S.L., Brunet A., Campisi J., Cuervo A.M., Epel E.S., Franceschi C., Lithgow G.J., Morimoto R.I., Pessin J.E., Rando T.A., Richardson A., Schadt E.E., Wyss-Coray T., Sierra F. (2014) Geroscience: linking aging to chronic disease. Cell. 2014 Nov 6;159(4):709-13. doi: 10.1016/j.cell.2014.10.039. PMID: 25417146; PMCID: PMC4852871. Elektronikusan elérhető itt: <https://pubmed.ncbi.nlm.nih.gov/25417146/>

López-Otín C., Blasco M.A., Partridge L., Serrano M., Kroemer G. (2023) Hallmarks of aging: An expanding universe. Cell. 2023 Jan 19;186(2):243-278. doi: 10.1016/j.cell.2022.11.001. Epub 2023 Jan 3. PMID: 36599349. Elektronikusan elérhető itt: <https://pubmed.ncbi.nlm.nih.gov/36599349/>

Moriya, Takafumi & Sato, Hiroshi & Ochiai, Tsubasa & Delcroix, Marc & Shinozaki, Takahiro (2023) Streaming End-to-End Target-Speaker Automatic Speech Recognition and Activity Detection. IEEE

Access. 11. 13906-13917. 10.1109/ACCESS.2023.3243690. Elektronikusan elérhető itt: https://www.researchgate.net/publication/368410409_Streaming_End-to-End_Target-Speaker_Automatic_Speech_Recognition_and_Activity_Detection

Nasreddine Z.S., Phillips N.A., Bédirian V., Charbonneau S., Whitehead V., Collin I., Cummings J.L., Chertkow H. (2005) The Montreal Cognitive Assessment, MoCA: a brief screening tool for mild cognitive impairment. J Am Geriatr Soc. 2005 Apr;53(4):695-9. doi: 10.1111/j.1532-5415.2005.53221.x. Erratum in: J Am Geriatr Soc. 2019 Sep;67(9):1991. doi: 10.1111/jgs.15925. PMID: 15817019. Elektronikusan elérhető itt: <https://pubmed.ncbi.nlm.nih.gov/15817019/>

Obádovics Cs., Tóth G. Cs. (2023) A magyarországi régiók népességének előreszámítása 2050-ig. KSH Statisztikai Szemle 2023. szeptember 20. Elektronikusan elérhető itt: <https://doi.org/10.20311/stat2023.09.hu0763>

OpenAI (2023). GPT-4 Technical Report. arXiv preprint arXiv:2303.08774. Elektronikusan elérhető itt: <https://arxiv.org/abs/2303.08774>

Radford, A., Kim, J.W., Xu, T., Brockman, G., Mcleavy, C. & Sutskever, I.. (2023). Robust Speech Recognition via Large-Scale Weak Supervision. <i>Proceedings of the 40th International Conference on Machine Learning</i>, in <i>Proceedings of Machine Learning Research</i> 202:28492-28518 Elektronikusan elérhető itt: <https://proceedings.mlr.press/v202/radford23a.html>

World Health Organization (WHO) (2021). Decade of Healthy Ageing: Baseline Report. Geneva: World Health Organization. Elektronikusan elérhető itt: <https://www.who.int/publications/i/item/9789240017900>