
STUDIA IURIS

JOGTUDOMÁNYI TANULMÁNYOK / JOURNAL OF LEGAL STUDIES

2025. II. ÉVFOLYAM 1. SZÁM



Károli Gáspár Református Egyetem
Állam- és Jogtudományi Doktori Iskola

A folyóirat a Károli Gáspár Református Egyetem Állam- és Jogtudományi Doktori Iskolájának a közleménye. A szerkesztőség célja, hogy fiatal kutatók számára színvonalas tanulmányaik megjelentetése céljából méltó fórumot biztosítson.

A folyóirat közlésre befogad tanulmányokat hazai és külföldi szerzőktől – magyar, angol és német nyelven. A tudományos tanulmányok mellett kritikus, önálló véleményeket is tartalmazó könyvismertetések és beszámolók is helyet kapnak a lapban.

A beérkezett tanulmányokat két bíráló lektorálja szakmailag. Az idegen nyelvű tanulmányokat anyanyelvi lektor is javítja, nyelvtani és stilisztikai szempontból.

A folyóirat online verziója szabadon letölthető (open access).

ALAPÍTÓ TAGOK

BODZÁSI BALÁZS, JAKAB ÉVA, TÓTH J. ZOLTÁN, TRÓCSÁNYI LÁSZLÓ

FŐSZERKESZTŐ

JAKAB ÉVA ÉS BODZÁSI BALÁZS

OLVASÓSZERKESZTŐ

GIOVANNINI MÁTÉ

SZERKESZTŐBIZOTTSÁG TAGJAI

BOÓC ÁDÁM (KRE), FINKENAUER, THOMAS (TÜBINGEN), GAGLIARDI, LORENZO (MILANO),
JAKAB ANDRÁS DSc (SALZBURG), SZABÓ MARCEL (PPKE), MARTENS, SEBASTIAN (PASSAU),
THÜR, GERHARD (AKADÉMIKUS, BÉCS), PAPP TEKLA (NKE), TÓTH J. ZOLTÁN (KRE),
VERESS EMŐD DSc (KOLOZSVÁR)

Kiadó: Károli Gáspár Református Egyetem Állam- és Jogtudományi Doktori Iskola

Székhely: 1042 Budapest, Viola utca 2-4

Felelős Kiadó: TÓTH J. ZOLTÁN

A tipográfia és a nyomdai előkészítés CSERNÁK KRISZTINA (L'Harmattan) munkája

A nyomdai munkákat a Robinco Kft. végezte, felelős vezető GEMBELA ZSOLT

Honlap: <https://ajk.kre.hu/index.php/jdi-kezdolap.html>

E-mail: doktori.ajk@kre.hu

ISSN 3057-9058 (Print)

ISSN 3057-9392 (Online)

URL: KRE ÁJK - Studia Iuris

<https://ajk.kre.hu/index.php/kiadvanyok/studia-iuris.html>

AZ INTELLIGENCIA HATÁRAI – A MESTERSÉGES INTELLIGENCIA SZABÁLYOZÁSA ÉS AZ ADATVÉDELEM ETIKAI LÉTJOGOSULTSÁGA

THE LIMITS OF INTELLIGENCE – THE REGULATION OF ARTIFICIAL INTELLIGENCE AND THE ETHICAL JUSTIFICATION OF DATA PROTECTION

FÜZESSÉRI FANNI¹

ABSTRACT ■ The rapid development of Artificial Intelligence (AI) raises numerous questions, particularly in the areas of regulation and data protection. This paper aims to present the different forms of AI (narrow, general, and superintelligent AI) and analyze the ethical and legal dilemmas that arise concerning data management and social impacts. It discusses the black box problem, the ethical issues surrounding the loan application dilemma, and the impact of data protection regulations such as GDPR and the AI Act on AI applications. The Clearview AI scandal and its data protection consequences highlight the growing intersection of AI and data protection, emphasizing the need for ethical AI development to prevent discriminatory effects and ensure individuals' rights in the digital age.

KEYWORDS: Artificial Intelligence (AI), data protection, ethics, GDPR, AI Act, black box problem, Clearview AI, data processing, discrimination, biometric data

ABSZTRAKT ■ A mesterséges intelligencia (MI) fejlődése számos új kérdést vet fel, különösen a szabályozás és adatvédelem terén. A cikk célja az MI különböző formáinak (szűk, általános, szuperintelligens MI) bemutatása, valamint az etikai és jogi dilemmák elemzése, amelyek az adatkezeléssel és a társadalmi hatásokkal kapcsolatosan merülnek fel. Részletesen tárgyalja a fekete doboz problémát, a hitelkérelem dilemma etikai kérdéseit, valamint az adatvédelmi jogszabályok, mint a GDPR és az AI Act hatásait az MI alkalmazásában. A Clearview AI botrány és annak adatvédelmi következményei kiemelik, hogy az MI és az adatvédelem egyre inkább összefonódik, és az etikus MI fejlesztése elengedhetetlen ahhoz, hogy elkerüljük a diszkriminatív hatásokat és biztosítsuk az egyének jogait a digitális korban.

KULCSSZAVAK: Mesterséges intelligencia (MI), adatvédelem, etika, GDPR, AI Act, fekete doboz probléma, Clearview AI, adatkezelés, diszkrimináció, biometrikus adatok

¹ PhD-hallgató, Károli Gáspár Református Egyetem, Állam- és Jogtudományi Doktori Iskola.

1. BEVEZETÉS

Azzal kevesen vitatkoznának, hogy a mesterséges intelligencia (a továbbiakban: MI) napjaink egyik legdinamikusabban fejlődő területe, amely alapjaiban változtatja meg mindennapjainkat és társadalmi rendszereinket. Azonban e technológia előrehaladtával olyan kérdések is felszínre kerülnek, amelyek a szabályozás és az adatvédelem etikai alapjaira helyezik a hangsúlyt. Ez az a terület, ahol a változás szinte percről percre megy végbe. Kezdjük az MI fogalmának tisztázásával, hiszen a kutatók jelenleg sem tudnak megegyezni a fogalmakról, illetve annyira gyorsan változik a technológia, hogy amit pár éve 'általános MI'-nek hívtunk, azt ma már 'gyenge MI'-nek tituláljuk.

2. AZ MI FOGALMA

Az MI-nek tehát nincs egységes definíciója, de általában olyan technológiát jelent, amely képes az emberi intelligenciához hasonló feladatok elvégzésére. NEIL WILKINS szerint azonban az MI-nek két különböző jelentése van a mindennapjainkban:

1. ANI ('Artificial Narrow Intelligence', magyarul 'szűk MI'): olyan összetett feladatokat képes végrehajtani, valamint olyan döntéseket tud hozni, amelyekre beprogramozták.²
2. AGI ('Artificial General Intelligence', magyarul 'általános MI'): emberi gondolkodásra képes gépek.³

Azonban Wilkins mellett MIHAIL CARADAICA is rámutat, hogy az AGI létezésére még nincs kézzelfogható bizonyítékunk. JOOTAEEK LEE⁴ szerint kiemelt figyelmet kellene fordítanunk az AGI-ra, hiszen várhatóan ezzel a fajta intelligenciával is számolnunk kell majd.

Ehhez a felsoroláshoz még egy harmadik elemmel is ki lehetne egészíteni a sort, mely szerint az ASI ('Artificial Superintelligence', magyarul 'szuperintelligens MI/ szuperintelligencia') eszébe jutna, hogy újratervezze saját magát, hogy még intelligensebbé váljon. IRVING JOHN GOOD (I. J. Good) 1965-ös hipotézise szerint

² MIHAIL CARADAICA PhD: *Artificial intelligence and inequality in European Union*. https://www.researchgate.net/publication/343859215_ARTIFICIAL_INTELLIGENCE_AND_INEQUALITY_IN_EUROPEAN_UNION.

³ NEIL WILKINS: *Artificial Intelligence: What You Need to Know About Machine Learning, Robotics, Deep Learning, Recommender Systems, Internet of Things, Neural Networks, Reinforcement Learning, and Our Future*. 2019. ISBN: 1795408561. E-book.

⁴ Jootaek Lee, a Rutgers Law School adjunktusa és jogi könyvtárosa.

„az ultraintelligens gép olyan gép, amely messze felülmúlja bármely ember bármilyen okos szellemi tevékenységét”.⁵ Ez alapján, ha egy MI elég intelligens ahhoz, hogy megértse saját tervezését, akkor újratervezheti önmagát, hogy 'még jobb' verzióját állítsa elő.⁶ Az ASI jelenleg nem elfogadott és nem bizonyított.

3. A BLACK BOX PROBLÉMA

Az emberek viselkedését és reakcióit az ismeretlentől való félelem már ősidők óta befolyásolja, és ez alól a mesterséges intelligenciától való félelem sem kivétel. Számos példa támasztja alá ezt az aggodalmat, és most a 'fekete doboz' (Black Box) problémát szeretném kiemelni, amely az átláthatóság hiányából ered. Ez azt jelenti, hogy az MI döntéshozatali folyamata sok esetben olyan bonyolult vagy megérthetetlen, hogy az emberek nem tudják követni, hogyan születnek a döntések, ami bizalomvesztéshez vezethet. A fekete doboz probléma különösen fontos az MI rendszerekben, mivel a mélytanulási algoritmusok sokszor úgy hozzák meg a döntéseket, hogy még a fejlesztők sem értik teljesen, hogyan zajlik a folyamat.⁷ A következő példában látni fogjuk, hogy például egy MI-alapú hitelbírálati rendszer megtagadhat egy kölcsönt anélkül, hogy egyértelmű indokot adna, ez pedig nemcsak jogi, hanem etikai problémákat is felvet ROBERT PASQUALE szerint. Ha az emberek nem értik, miért született egy adott döntés, akkor hogyan lehet felelősségre vonni az MI-t vagy a mögötte álló szervezeteket?

4. A HITELKÉRELEM DILEMMA PROBLÉMÁJA

Képzeli el, hogy egyre több bank fogja az MI-t használni a hitelkérelmek elbírálásánál is, hiszen az MI a legpontosabban, leggyorsabban és leghatékonyabban tudja értékelni az adós helyzetét és a bank számára a legkedvezőbb döntést tudja hozni, amely szinte minden esetben biztosítja a hitel visszafizetésének magas százaléku esélyét. Egy elutasított kérelmező azonban azt veszi észre, hogy vagy faji vagy valamilyen más okból el lett utasítva nemcsak az ő kérelme, hanem a

⁵ IRVING JOHN GOOD: *Speculations Concerning the First Ultraintelligent Machine*. *Advances in Computers*. <https://www.sciencedirect.com/science/article/abs/pii/S0065245808604180> (Letöltve: 2025.02.11.).

⁶ STEPHEN MCALEESE: *How Do AI Timelines Affect Existential Risk?* Cornell University, 2022. <https://arxiv.org/pdf/2209.05459.pdf>.

⁷ FRANS L. LEEUW: *Evaluating the use of artificial intelligence and big data in policy making. Unpacking black boxes and testing white boxes*. 2024, 84.

többieké is. Ezért pert indít, melyre a bank azt válaszolja, hogy az algoritmus direkt nem veszi figyelembe a kérelmező faji vagy bármilyen más hovatarozását. Ennek ellenére a statisztikák mást mutatnak.⁸

Összegezve rájövünk, hogy az algoritmus valójában csak a dolgát végezte, hiszen az lett neki megadva vagy a bank számára a lehető legjobb kérelmezőket válassza. Azonban kiderül, hogy a 'jó' kérelmezők jómódú körülmények között élnek vagy éltek, és az algoritmus azokat utasította el, akik szegény környéken élnek vagy nőttek fel. Sok esetben ezek akár feketék is lehettek.⁹ Ahhoz, hogy megelőzzük az ilyesfajta diszkriminációt, olyan MI algoritmusok fejlesztésére kell törekednünk, amelyek átláthatóak az ellenőrzés számára, és skálázhatóak.¹⁰

Azonban, ha jobban belegondolunk, lehetséges, hogy nem sikerül majd az MI-vel megoldanunk ezt a problémát, hiszen abban a kérdésben, mely szerint a gazdagabbak a 'jó' kérelmezők, akiknek érdemes hitelt nyújtani, az MI nem fog tudni jó döntést hozni. Itt valójában egy sokkal mélyebb, rendszerszintű problémáról van szó. Ahogyan láthatjuk, az MI 'csak végrehajt'.

Az egyik jó megoldás lehet erre a problémára, hogy 'külön algoritmus' kezeli az X bevételi összeg feletti és alatti kérelmezők kérelmeit. Esetleg lehetne külön algoritmus a feketék és a fehérek kérelmeire. Ez viszont csak abban az esetben lehetne jobb döntés, ha a mai helyzetet vennénk figyelembe. Az elkövetkezendő évtizedekben azonban remélhetőleg szűkülni fog a gazdasági különbség a feketék és a fehérek (valamint gazdagok és szegények) között.

NICK BOSTROM és ELIEZER YUDKOWSKY, ezen dilemma megalkotói szerint nem örülnénk, ha a világon egyik bank sem engedne felvenni hitelt, és amikor megkérdeznénk a bankot, hogy indokolja meg, miért nem kaptunk, nem tudnák megmondani, és ki sem lehetne deríteni, hogy pontosan miért döntött így a gép.

Ez a dilemma nem csupán etikai kérdéseket vet fel, hanem a jogalkotók számára is fontos jogi problémákat generál. Hogyan szabályozzuk az MI algoritmusokat úgy, hogy azok ne hozzanak diszkriminatív döntéseket, ugyanakkor ne korlátozzák az MI által nyújtott hatékonyságot és gyorsaságot? Hogyan biztosítható, hogy az algoritmusok átláthatóak és nyomon követhetők legyenek, miközben megfelelnek az adatvédelmi jogoknak és a társadalmi igazságosság elveinek?

Ezzel összefüggésben az emberek szeretnék megelőzni a lehetséges kockázatok, így megjelent az igény az MI-vel összefüggésben az adatok szabályozására törvényi szinten is a világ minden részén. Habár az MI rendszerek a gazdaságtól az egészségügyig, az oktatástól a közlekedésig számos területen kínálnak

⁸ NICK BOSTROM – ELIEZER YUDKOWSKY: *The Ethics of Artificial Intelligence*. 2014.

⁹ KEITH FRANKISH – WILLIAM M. RAMSEY: *The Cambridge Handbook of Artificial Intelligence*. Cambridge University Press, 2014, 316–334.

¹⁰ Ibid.

megoldásokat, ennek ellenére olyan társadalmi kihívásokkal is szembesülünk nap, mint nap, mint az adatokkal való visszaélések, az átláthatóság hiánya és az automatizációval járó munkaerőpiaci átalakulás.

5. ADATVÉDELMI JOGSZABÁLYOK, MINT A GDPR, AZ AI ACT ÉS AZ INFÓTÖRVÉNY

Az információs önrendelkezési jogról és az információszabadságról szóló törvény [2011. évi CXII. törvény] (a továbbiakban: Infótörvény) mellett az Európai Unió két meghatározó jogszabálya, a GDPR [Az Európai Parlament és a Tanács 2016. április 27-i (EU) 2016/679 rendelet (általános adatvédelmi rendelet)] és az AI Act [Az Európai Parlament és a Tanács 2024. június 13-i (EU) 2024/1689 rendelet (a mesterséges intelligenciáról szóló rendelet)] azok, amelyek a digitális térben történő adatkezelést és a mesterséges intelligencia alkalmazását szabályozzák. A GDPR 2018. május 25-én lépett hatályba, célja pedig a személyes adatok védelme és az adatok szabad áramlásának biztosítása. Bár az AI Act 2024. augusztus 1-jén hatályba lépett, rendelkezéseinek gyakorlati alkalmazása több lépésben valósul meg. Az I. és II. fejezet – amelyek az általános rendelkezéseket, illetve a tiltott MI-gyakorlatokat szabályozzák – már 2025. február 2-től alkalmazandó. Ezt követően, 2025. augusztus 2-től lépnek hatályba a III. fejezet 4. szakaszának (a bejelentő hatóságokra és bejelentett szervezetekre vonatkozó) előírásai, valamint az általános célú MI-modellekre (V. fejezet), az irányítási struktúrákra (VII. fejezet), a szankciókra (XII. fejezet, kivéve a 101. cikket), és a titoktartásra (78. cikk) vonatkozó rendelkezések. A rendelet teljes körű alkalmazása azonban csak 2026. augusztus 2-án kezdődik meg, míg a magas kockázatú MI-rendszerek besorolására irányadó 6. cikk (1) bekezdése, valamint a korábban forgalomba hozott rendszerekre vonatkozó megfelelési követelmények végső határideje 2027. augusztus 2.¹¹ Célja pedig az MI rendszerek biztonságos és etikus használatának biztosítása. Az Infótörvény már 2012 óta szabályozza Magyarországon az adatvédelemmel összefüggő jogsértéseket, azonban a 2018. évi módosítás óta már a GDPR-al összhangban.

A GDPR minden olyan szervezetre vonatkozik, amely az EU-ban élő személyek adatait kezeli, függetlenül attól, hogy a szervezet földrajzilag hol található. Az adatkezelésnek jogszerűnek kell lennie, például az érintett hozzájárulása vagy szerződés teljesítése alapján. Fontos elvek a GDPR-ban az átláthatóság, az

¹¹ European Parliament – AI Act implementation timeline. [https://www.europarl.europa.eu/RegData/etudes/ATAG/2025/772906/EPRS_ATA\(2025\)772906_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/ATAG/2025/772906/EPRS_ATA(2025)772906_EN.pdf).

adatbiztonság és az érintettek jogainak biztosítása. Az adatkezelők kötelesek megfelelő technikai és szervezési intézkedéseket bevezetni, hogy megvédjék az adatokat. Ha egy szervezet megsérti a rendeletet, akár 20 millió eurós vagy a globális árbevétel 4%-ának megfelelő bírságot is kaphat.

Az AI Act az MI-rendszereket kockázati szintek alapján kategorizálja. Tiltottak azok a rendszerek, amelyek elfogadhatatlan kockázatot jelentenek, például társadalmi pontozásra használt MI-rendszerek. A magas kockázatú rendszerek szigorú szabályozás alá esnek, például az önéletrajzszűrő algoritmusok. A korlátozott kockázatú rendszereknek átláthatósági követelményeknek kell megfelelniük, például chatbotok esetében jelezni kell, hogy a felhasználó MI-vel kommunikál. A minimális kockázatú MI-k, például videojátékokban vagy spam-szűrőkben alkalmazott rendszerek külön szabályozás nélkül is működhetnek.

Mindkét jogszabály célja az egyének jogainak védelme a digitális térben. A GDPR elsősorban a személyes adatok kezelésére, míg az AI Act az MI rendszerek fejlesztésére és működtetésére összpontosít. Az MI fejlesztőinek és üzemeltetőinek mindkét rendelet követelményeit figyelembe kell venniük, hiszen egy MI rendszer gyakran személyes adatokat is kezel. Az AI Act által tiltott MI gyakorlatok megsértése esetén a legnagyobb, akár 35 millió eurós vagy a globális árbevétel 7%-ának megfelelő bírság is kiszabható.

Ezek a szabályozások biztosítják, hogy az Európai Unióban a technológiai fejlődés az egyének jogainak és biztonságának tiszteletben tartásával történjen. A GDPR és az AI Act együttes alkalmazásának a célja, hogy a digitális gazdaságban az adatkezelés és a mesterséges intelligencia használata átlátható, biztonságos és etikus maradjon.

A mesterséges intelligencia és az adatvédelem összefonódása számos adatvédelmi incidenst eredményezett, egy példán keresztül szeretném bemutatni az adatvédelem fontosabb fogalmait.

6. A CLEARVIEW AI BOTRÁNY

A Clearview AI nevű amerikai cég 2020-ban egy olyan arcfelismerő technológiát fejlesztett, amely nyilvánosan elérhető képek és videók alapján képes azonosítani személyeket. A vállalat több milliárd képet gyűjtött össze anélkül, hogy az érintettek hozzájárulását kérte volna, és ezeket az adatokat különböző szervezetek számára értékesítette.¹²

¹² JULIET SAMANDARI – AMELIA SAMANDARI: *Biometric Protection and Security: A Case Study on Clearview AI*. 2024, 278–283.

A vállalat által gyűjtött képek és az azokhoz kapcsolódó biometrikus adatok személyes adatnak minősülnek, mivel ezek alapján az egyének egyértelműen azonosíthatók. Az adatkezelést a Clearview AI végezte azáltal, hogy ezeket az adatokat gyűjtötte, tárolta és feldolgozta. Mivel a vállalat maga döntött az adatkezelés céljairól és eszközeiről, így adatkezelőnek minősült, de nem adatfeldolgozónak, mivel nem más szervezetek megbízásából kezelte az adatokat.

Az érintettek azok az egyének voltak, akiknek a képeit és biometrikus adatait a Clearview AI gyűjtötte, azonban a vállalat nem rendelkezett jogalappal ezen adatok kezelésére, mivel nem kérte az érintettek hozzájárulását. Bár nem történt konkrét adatvédelmi incidens, a Clearview AI tevékenysége sértette az érintettek jogait és a GDPR előírásait. A vállalat nem nevezett ki adatvédelmi tisztviselőt ('Data protection officer', röviden: DPO), noha erre a nagymértékű megfigyelés miatt kötelezettsége lett volna a 37. cikk szerint.

Továbbá az érintettek nem tudtak élni jogaikkal, mivel nem kaptak tájékoztatást arról, hogy adataikat gyűjtik és használják. A Clearview AI nem végzett adatvédelmi hatásvizsgálatot sem, ami azt jelentette, hogy nem mérte fel előzetesen az adatkezeléssel járó kockázatokat. Ezek a hiányosságok jelentős GDPR-sértéseknek minősültek, ami miatt a vállalatot Franciaország és Olaszország adatvédelmi hatósága is elmarasztalta és a kiszabható maximális összeggel, 20 millió eurós pénzbírsággal sújtotta.

7. AZ ADATVÉDELEM ETIKAI JELENTŐSÉGE

A Clearview AI esete jól szemlélteti, hogyan sértheti meg egy AI-alapú szolgáltatás a GDPR előírásait, különösen a személyes adatok gyűjtésére, kezelésére és felhasználására vonatkozóan.

Az adatvédelem etikai létjogosultsága abból fakad, hogy a személyes adatok védelme az egyén méltóságának, autonómiájának és szabadságának alapvető feltétele. Az Európai Unió Alapjogi Chartája [Az Európai Unió Alapjogi Chartája (2016/C 202/02)] és a GDPR egyaránt kimondja, hogy az egyéneknek joguk van személyes adataik védelméhez, és csak meghatározott jogalapok alapján lehet azokat kezelni. A MI fejlődése új kihívások elé állítja ezt az alapelvet, mivel az MI rendszerek gyakran nagy mennyiségű személyes adatot dolgoznak fel az érintettek tudta vagy beleegyezése nélkül.

Az adatvédelem etikai jelentősége túlmutat a jogi előírásokon, mivel a technológiai fejlődés nem mindig halad egyenes arányban az emberi jogok védelmével. Az 'etikai MI' koncepciója szerint az MI fejlesztése és alkalmazása során figyelembe kell venni az egyének jogait, biztosítva a transzparenciát, a tisztességes

adatkezelést és a döntéshozatal igazságosságát. Ez különösen fontos az automatizált döntéshozatali rendszerek esetében, amelyek potenciálisan diszkriminatív hatással lehetnek bizonyos társadalmi csoportokra, ha az alapjukat képező adatok elfogultak vagy hiányosak (ahogyan a hitelkérelem dilemma problematikájánál látható volt).

Az adatvédelmi etika egyik kulcskérdése az érintettek beleegyezésének problémája. Az MI által feldolgozott adatok gyakran rejtett vagy indirekt módon kerülnek gyűjtésre, például közösségi média platformokon vagy nyilvánosan elérhető forrásokból, ahogyan azt a Clearview AI esete is megmutatta. Luciano Floridi szerint az ilyen gyakorlatok azt a kérdést vetik fel, hogy a felhasználók valóban informált döntést hoznak-e az adataik megosztásáról, vagy csupán passzív adatforrásként kezelik őket. Az alacsony adatvédelmi tudatosság és a bonyolult felhasználási feltételek miatt sokan nincsenek tisztában azzal, hogy milyen mértékben használják fel adataikat, ami sérti az autonómiájukat és önrendelkezési jogukat.

Az adatvédelem nemcsak jogi, hanem erkölcsi szükségszerűség is, amely biztosítja, hogy az MI fejlődése ne veszélyeztesse az egyéni szabadságjogokat, hanem inkább azok tiszteletben tartását és megerősítését szolgálja. Az elmúlt évek gazdasági fejlődésének fényében indokolt a harmadik országok gazdasági szereplőivel folytatott együttműködés alapos vizsgálata, mivel ezen országokban a GDPR és az AI Act szabályozása jelentős hiányosságokat mutat. Az elkövetkezendő évek fontos feladata lesz a jogharmonizáció megteremtése a harmadik országbeli szereplőkkel.

A mesterséges intelligencia fejlődése izgalmas és kihívásokkal teli korszakot nyitott meg. Az intelligencia határai nemcsak technológiai kérdések, hanem mély etikai és társadalmi dilemmák is. Az MI szabályozása és az adatvédelem garantálása nem csupán indokolt, hanem alapvetően szükséges annak érdekében, hogy a technológia valóban az emberiség érdekét szolgálja.

IRODALOMJEGYZÉK

NICK BOSTROM – ELIEZER YUDKOWSKY: *The Ethics of Artificial Intelligence*. 2014.

MIHAIL CARADAICA PhD: *Artificial intelligence and inequality in European Union*. https://www.researchgate.net/publication/343859215_ARTIFICIAL_INTELLIGENCE_AND_INEQUALITY_IN_EUROPEAN_UNION

KEITH FRANKISH – WILLIAM M. RAMSEY: *The Cambridge Handbook of Artificial Intelligence*. Cambridge University Press, 2014

IRVING JOHN GOOD: *Speculations Concerning the First Ultraintelligent Machine*. *Advances in Computers*, <https://www.sciencedirect.com/science/article/abs/pii/S0065245808604180>

FRANS L. LEEUW: *Evaluating the use of artificial intelligence and big data in policy making. Unpacking black boxes and testing white boxes.* 2024

STEPHEN McALEESE: *How Do AI Timelines Affect Existential Risk?* Cornell University, 2022. <https://arxiv.org/pdf/2209.05459.pdf>

JULIET SAMANDARI – AMELIA SAMANDARI: *Biometric Protection and Security: A Case Study on Clearview AI.* 2024

NEIL WILKINS: *Artificial Intelligence: What You Need to Know About Machine Learning, Robotics, Deep Learning, Recommender Systems, Internet of Things, Neural Networks, Reinforcement Learning, and Our Future.* 2019. ISBN: 1795408561. E-book