

Pécsi Tudományegyetem
Állam- és Jogtudományi Doktori Iskola

Diósi Szabolcs

Mesterséges intelligencia, szintetikus valóság

- **Az MI és GenMI rendszerekkel kapcsolatos globális kihívások
és európai szabályozási stratégiák**

Doktori (PhD) értekezés



Témavezető:

Dr. Barcsi Tamás

Pécs

2024

„A fejlődés ellen nincs gyógymód”

Neumann János

Műhelyvitára benyújtott, nem véglegesített változat.

1	Bevezetés.....	5
1.1	A disszertáció kutatási témája és a témaválasztás indokolása	5
1.2	A disszertáció tartalmi áttekintése	7
1.3	Módszertan, a kutatás hasznosítása	9
1.4	A dolgozat központi kérdései, hipotézisei.....	10
2	Az MI (rövid) története	13
2.1	Theseus-tól a ChatGPT-ig.....	13
2.2	Az MI-ről általában	18
2.3	Szűk és Általános MI.....	19
2.4	Az öntanuló és önfejlesztő MI	21
3	Alkalmazások és kihívások.....	26
3.1	MI a köz szolgálatában?	26
3.1.1	RPA és ADM - az automatizált döntéshozó rendszerek	26
3.2	Prediktív analitika az adatvezérelt döntéshozatal szolgálatában	28
3.2.1	Adózással és szociális juttatásokkal kapcsolatos csalások	28
3.2.2	Rendvédelmi szervek és büntető igazságszolgáltatás	29
3.3	Az MI alkalmazásból eredő kihívások, kockázatok	32
3.3.1	A Black Box probléma – átláthatóság, értelmezhetőség, indoklási kötelezettség	32
3.3.2	Egyenlő bánásmód követelményének megsértése	35
3.3.3	Elfogult és diszkriminatív döntéshozatal a gyakorlatban.....	36
3.3.4	Szabadság, demokrácia és a döntési autonómia	39
3.3.5	Kinek a választása?.....	40
3.4	Előre kalkulált kockázat, felülről szabályozott felhasználás	43
4	Az Unió MI-rendelete	46
4.1	Európa az MI szabályozás útján.....	46
4.2	A kockázat alapú megközelítés	48
4.3	Korlátozott és minimális kockázatú MI-rendszerek	49
4.4	Nagy kockázatú MI-rendszerek.....	50
4.5	Tiltott gyakorlatok, „elfogadhatatlan kockázat”	52
4.5.1	Magatartást torzító gyakorlatok.....	52
4.5.2	Általános társadalmi pontozás	54
4.5.3	Biometrikus azonosítási rendszerek valós idejű használata	56
4.6	Az Európai Unió Tanácsának közös álláspontja.....	57
4.6.1	Új kategória: Általános célú MI	60
4.7	A Parlament kompromisszumos javaslata.....	61
4.7.1	Széles körű alkalmazás, általános célú felhasználás	64
4.7.2	Az alapmodellek szolgáltatóinak kötelezettségei és a ChatGPT szabály	65
4.8	Háromoldalú tárgylások.....	67

4.9	A 2023 március 13-án elfogadott törvényszöveg	69
4.9.1	Mi az MI?	69
4.9.2	Kockázatos gyakorlatok	69
4.10	Két szintű szabályozási megközelítés	76
5	Szintetikus valóság	81
5.1	Az új technológia: GenMI	81
5.2	LLM – új korszak a digitális tartalomgyártásban?	82
5.3	LLM, mint sztochasztikus papagáj	83
5.4	Megbízhatatlan LLM – Hallucináció, hamis információ és szándékos félrevezetés	85
5.4.1	Megjelenési formái, típusai és lehetséges okai	86
5.4.2	Hallucináció orvoslása	87
5.4.3	Szándékos félrevezetés szimulált környezetben	88
5.5	A modell képzésből és tanításból eredő kihívások	91
5.5.1	Kinek a tanító adata?	92
5.5.2	Modell összeomlás	95
5.6	Szintetikus vagy valódi szöveges tartalom?	97
5.7	(Gen)MI-t hoz a jövő?	98
6	Szép új (online) világ	103
6.1	Zajos terek, kétes információk	104
6.2	A figyelmed eladó	107
6.3	A jövőd is eladó?	108
6.4	Információs buborékok, személyre szabott valóságok	109
6.4.1	Az Információfeldolgozás torzulása	110
6.5	Az igazság hanyatlása és a tényeken túli valóság	113
6.6	Félretájékoztatás, dezinformáció	115
6.7	A dezinformáció szolgálatában - Mikro-célzás, állhírek és botok	117
6.8	GenMI a dezinformáció szolgálatában	120
6.9	Szintetikus hang- és videótartalmak	121
6.10	Az Európai Unió dezinformációs stratégiája	128
6.11	Digitális Szolgáltatásokról Szóló Rendelet	130
6.11.1	Átláthatósági és tájékoztatási kötelezettség	131
6.11.2	A személyre szabott célzott hirdetés és manipuláció tilalma	132
6.12	Észrevételek	134
6.13	Irodalomjegyzék	138
6.14	Publikációs jegyzék	161

1 Bevezetés

1.1 A disszertáció kutatási témája és a témaválasztás indokolása

A dolgozat a Mesterséges Intelligencia és a Generatív Mesterséges Intelligencia technológiák közelmúltban tapasztalt látványos fejlődésének és széleskörű elterjedésének társadalmi következményeit, valamint az azokra adott Uniós jogalkotási stratégiákat vizsgálja. A téma kétségtelenül népszerű, az elmúlt években számos tudományos igényességű monográfia, könyv, tanulmány és cikk született a kérdéskörrel.¹ Aktualitását jelzi az is, hogy jelen dolgozat írásával párhuzamosan alkotta meg az Európai Unió azon MI rendeletét, amely a világon elsőként vállalkozik egy olyan átfogó jogi keret létrehozása, amely garantálja az ilyen-rendszerek fejlesztésének, használatának és forgalmazásának biztonságos feltételeit az EU határain belül.² A szabályalkotási folyamat során az Unió egy olyan jogi környezet kialakítására törekedett, amely nyitott a technológiai újítások felé, ösztönzi a kutatást és fejlesztést, egyúttal képes minimalizálni a társadalmi kockázatokat. Az innováció és európai versenyképesség megőrzésének támogatása, illetve az alapvető emberi jogok és európai értékek következetes, kompromisszumot nem ismerő védelme, mint olykor egymást korlátozó célok feloldására, pedig kockázatalapú szabályozási keretet javasolt, amely az MI rendszerek által jelentett kockázat mértékével arányos kötelezettségeket ír elő.

¹ A technológia növekvő befolyásának és szüntelen bővülő alkalmazási lehetőségeinek köszönhetően a témakör a szélesebb közvélemény érdeklődésének is a középpontjába került. A lehetséges felhasználási módzatokra – közel sem kimerítő - példaként szolgálhat, hogy már napjainkban is sok helyen MI-alapú rendszerek automatizálják az ügyintézt és az adminisztrációt, Nagy Nyelvi Modellekre épülő virtuális asszisztensek javítják az ügyfélszolgálat minőségét, az adóhatóságok MI-t alkalmaznak az adócsalás és adóelkerülés felderítésére, a szociális ellátórendszerek az állami segítségnyújtási ellátásokra és szolgáltatásokra való jogosultságának értékelésére. A rendvédelem és igazságszolgáltatás területein MI kockázatértékelő algoritmusok elemzik a bűnügyi adatokat, előrejelzik a bűncselekmények elkövetésének valószínűségét, támogatják a nyomozást és a bírói döntéshozatalt. Biometrikus arcfelismerőrendszerek segítik a bűnmegelőzést és a határellenőrzést. Az egészségügy területén MI-alapú rendszerek elemzik a radiológiai felvételeket, azonosítják a gyanús elváltozásokat és betegségekre utaló jeleket. A precíziós mezőgazdaságban MI-vezérelt drónok, robotok és szenzorrendszerek monitorozzák a növények egészségi állapotát, azonosítják a kártevőket, betegségeket, valamint optimalizálják a vízfelhasználást és egyéb erőforrásokat. A környezetvédelem területén MI-modellek elemzik az ökológiai adatokat, előrejelzik a klímaváltozás hatásait és támogatják a fenntartható erőforrás-gazdálkodást. Bankok és biztosítótársaságok algoritmusokat használnak az ügyfelek hitelképességének értékelésére, hitelkérelmek gyorsabb feldolgozására a kockázatok csökkentésére. A humán erőforrás-menedzsment területén az MI alkalmazható a jelentkezők toborzására, kiválasztására, üres álláshelyek meghirdetésére, pályázatok szűrésére, valamint a jelöltek interjúk során történő értékelésére. Az e-kereskedelemben és a szórakoztatóiparban egyaránt intelligens ajánlórendszerek elemzik a fogyasztói preferenciákat, hogy célzott, személyre szabott termékeket és tartalmakat kínáljanak.

² Az Európai Unió Tanácsa 2024. május 21-én hagyta jóvá a végleges törvényszöveget, melyet várhatóan 2024 júliusában fognak közzétenni az Európai Unió Hivatalos Lapjában.

A kezdeti a jogalkotási folyamat irányát változtatta meg a GenMI technológiák széles körben történő elterjedése és szabadon hozzáférhetővé válása. Hosszú ideje ismert kihívás a technológiai ágazat jogalkotási eljárásai kapcsán, hogy az innováció hajlamos gyorsabban haladni, mint a szabályozás, mire a szabályozó hatóságok feltérképezik az új technológiákat és kidolgozzák a megfelelő keretrendszert, addigra újabb fejlesztések jelennek meg. Az MI-hez hasonló felforgató technológiák szabályozásakor kiváltképp igaz Koltay megállapítása miszerint, *a jogi szabályozás rendszerint követő üzemmódban van.*³ Az idő szorításában születő jogalkotói stratégiákat fenyegető további veszély, hogy a döntéshozók nem rendelkeznek kellő mennyiségű információval, vagy ha igen azokból nem a legmegfelelőbbeket választják ki a szabályozás megalapozásához. Ilyen helyzetekben a jogalkotók könnyen kerülhetnek abba a dilemmába, hogy a *"meggondolatlan cselekvés és a teljes bénultság"* között kell választaniuk.⁴

Az Európai Bizottság által 2021-ben közzétett rendelettervezetben foglalt szabályozási paradigma a GenMI technológiák robbanásszerű elterjedését megelőzően lett megfogalmazva, ebből kifolyólag nem volt maradéktalanul alkalmas az újonnan megjelenő kihívások kezelésére. A kockázatalapú megközelítés a rendszerek előre beazonosított alkalmazási területeire és felhasználási módozataira összpontosít (például az olyan kiemelten veszélyes területekre, mint a bűnüldözés, a határvédelem, vagy a kritikus infrastruktúrák működtetése). A GenMI technológiák esetén az ilyen típusú szabályozások bár bizonyos mértékben hasznosak lehetnek nem képesek lefedni az összes potenciális veszélyforrást. A dolgozat első részének középpontjában ezért az EU MI jogszabályalkotási folyamatának feltérképezése áll. A dolgozat röviden bemutatja az MI rendszerek felhasználási területeit és az ezekből adódó kockázatokat, majd részletesen elemzi a szabályozás kialakításának folyamatát, külön kiemelve, hogy a GenMI modellek megjelenése milyen hatással volt az EU kezdeti, kockázatalapú megközelítésére, illetve hogyan miként kényszerítette a jogalkotókat eredeti szabályozási stratégiájuk átgondolására.

A dolgozat második része a GenMI technológiák elterjedéséből eredő kihívásokat veszi górcső alá. Többek között kitér ChatGPT-hez hasonló Nagy Nyelvi Modellek működési sajátosságaiból eredő olyan releváns kockázatokra mint a gépi hallucináció, a szándékos megtévesztés, a szintetikus adatokon történő modellképzés, az információs homogenizálódás

³ Koltay András (2021): Előszó. In Török Bernát és Zódi Zsolt (szerk) A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 11.o.

⁴ G. Karácsony Gergely (2020): Okoseszközök – Okos jog? A mesterséges intelligencia szabályozási kérdései. Budapest, Dialóg Campus. 77.o.

és a modell összeomlás kérdéskörei. Vizsgálódásának fókuszát a 21. századi információs környezet fokozatos eltorzulásához vezető több évtizede tartó folyamat feltérképezésére helyezi. A dolgozat feltételezése szerint az online terek információs közegét súlytó torzító folyamatokat a szöveges és audiovizuális szintetikus tartalmakat megközelítőleg korlátlan mennyiségben előállítani képes GenMI modellek robbanásszerű elterjedése jelentős mértékben felgyorsítja majd.

1.2 A disszertáció tartalmi áttekintése

Az első fejezetben az MI fejlődéstörténetének rövid bemutatására és a technológiacsáláddal kapcsolatos alapfogalmak ismertetésére kerül sor. Külön figyelem összpontosul a modern értelemben vett MI-t jellemző gépi tanulási mintázatok feltérképezésére. Ennek jelentősége nem elhanyagolható, hiszen míg korábban a hagyományos számítógépes rendszerek működése mereven kötött volt a programozók által előre megírt utasításokhoz, melyek. lépésről lépésre határozták meg, hogy pontosan mit kell tennie a rendszereknek egy adott feladat végrehajtásához, a gépi tanulással lehetővé vált az MI számára, hogy az adatminták összefüggéseire és korábbi tapasztalataira alapozva önállóan (iteratív módon) tanuljon és fejlődön, anélkül, hogy további instrukciókra lenne szüksége.

A második fejezet különböző esettanulmányokon keresztül mutat be olyan napjainkban is elterjedt – elsősorban a közszférában használt– MI alkalmazásokat melyek érdemi hatással lehetnek az állampolgárok jogaira és kötelezettségeire (pl.: bíróságok, adóhatóságok, bűnüldöző szervek által alkalmazott rendszerek). A gyakorlatok kiválasztása nem önkényes alapon történt, segítségükkel könnyebben megérthető, hogy a 2016-ban kezdődő EU-s jogalkotási stratégiák milyen veszélyeket és kockázatokat azonosítottak e rendszerekkel kapcsolatban. A fejezet második részében ismertetésre kerülnek a gyakorlati alkalmazásból eredő lejelentősebb szabályozási és állampolgári kihívások, nevezetesen: (1) a gépi tanulással képzett MI-működésére jellemző átláthatóság, elszámoltathatóság és kiszámíthatóság hiánya; (2) az MI által vezérelt téves, elfogult vagy diszkriminatív döntéshozatal veszélye; (3) az MI-rendszerek alapvető emberi szabadságjogokra, döntési autonómiára és a demokratikus intézményekre gyakorolt káros hatása.

A harmadik fejezetben az MI-rendelet jogalkotási folyamatának bemutatására kerül sor, külön hangsúlyt fektetve a nagy és elfogadhatatlan kockázatúkat kategorizált MI alkalmazások

területén eszközölt jogszabályi változásokra, és az *általános célú MI modelleket* érintő új szabályozási keretek kialakítására. A fejezet időrendben végig vezeti az olvasót az MI-rendelet megalkotásának folyamatán, kezdve a Bizottság 2021. április 21-én közzétett rendelet-tervezetével, a Tanács 2022. decemberi közös álláspontjával és a Parlament 2023. májusi kompromisszumos javaslatával. Ezt követően kitér a 2023 második felében zajló háromoldalú egyeztetésekre, valamint a jogszabály 2024. március 13-ai végleges elfogadásának körülményeire is.

A negyedik fejezet a szélesebb köztudatba 2022 végén berobbanó GenMI modellek sajátosságainak feltérképezésére vállalkozik. E rendszerek lényege abban rejlik, hogy felhasználói utasítások (promptok) alapján új tartalom létrehozására képesek, és hogy eredeti rendeltetésük mellett számos további alkalmazási lehetőséggel is bírnak. Míg a korai modellek csupán egyetlen modalitás, azaz egyetlen típusú tartalom előállítására korlátozódtak (például a Nagy Nyelvi Modellek csak szöveges bementet tudtak értelmezni és szöveges tartalmat tudtak előállítani), a legújabb fejlesztésű multimodális modellek már több eltérő formátumú tartalom feldolgozására és előállítására is alkalmasak. Az ilyen multimodális GenMI-k egyik látványos példája a 2024. júniusában minden felhasználó számára szabadon elérhetővé tett Luma Dream Machine modellje, amely kép vagy szöveges bemenetek alapján körülbelül két perc alatt képes 15 másodperces hosszúságú professzionális filmstúdiók termékeit idéző videók és animációk előállítására.⁵ A fejezet emellett kitér a ChatGPT-hez hasonló Nagy Nyelvi Modellek széles körű alkalmazásából eredő olyan releváns kihívásokra, mint a gépi hallucináció, a szándékos megtévesztés, a szintetikus adatokon történő modellképzés, az információs homogenizálódás és a modell összeomlás kérdéskörei.

Az ötödik fejezet első fele a GenMI modellek által előállított szintetikus tartalmak elterjedésének hosszútávú kockázatait járja körbe. Annak érdekében, hogy a folyamat teljes egészében feltérképezhetővé váljon a fejezet egészen a 2000-es évek elejéig – a Web2.0 kialakulásáig - tekint vissza az időben, és e ponttól kezdődően mutatja be, hogy miként változott meg a nyilvánosság szerkezete az elmúlt három évtizedben, valamint, hogy az algoritmus eredetű információtorzítás milyen szerepet játszott a társadalmi és politikai polarizáció kiélesítésével, az állhírek és a dezinformáció elterjedésével, valamint az *igazság utáni korszak*

⁵ A transzformer architektúrán és GAN gépi tanuláson alapuló GenMI segítségével bármely regisztrált felhasználónak havonta harminc ingyenes videó elkészítésére van lehetősége. Előfizetés esetén ez a tartalomgyártási kvóta számottevően növekszik kitolódik. A szoftver elérhető: <https://lumalabs.ai/dream-machine>

eljövetelével összefüggésben. A fejezet többek között kitér a figyelemgazdaság, a profilalapú tartalomszűrő algoritmusok, a visszhangkamrák és *kiberbalkanizáció* jelenségeire, de ezzel párhuzamosan az emberi információfeldolgozásra és valóságértelmezésre jellemző kognitív információs torzítások rövid bemutatását is megkísérli. A dolgozat ezen része emellett, azzal a kérdéssel is foglalkozik, hogy napjaink több tekintetben is kiszolgáltatott információs közegében milyen következményei lehetnek annak, ha a szintetikus tartalmakat korlátlan mennyiségben, gyorsan, olcsón előállítani képes GenMI modellek minden felhasználó számára egyetlen kattintás távolságba kerülnek.

A fejezet második fele az EU digitális szolgáltatásokról szóló rendeletének (DSA) ismertetésre tesz kísérletet. A jogszabály – melynek célja, hogy biztonságos, kiszámítható és megbízható online környezetet teremtsen az Unió összes állampolgára számára – szorosan kötődik a dolgozatban tárgyalt kihívásokhoz, minthogy az MI-rendelet és a GDPR mellett ebben a rendeletben fogalmazták meg azokat a szolgáltatói kötelezettségeket, melyek az automatikus gépi döntések és az algoritmikus tartalomszűrés átláthatóságának garantálására, valamint az egyén döntési szabadságát befolyásoló manipulatív gyakorlatok korlátozására irányulnak. A fejezet arra a következtetésre jut, hogy bár a DSA építő és szükséges kiegészítője azon már hatályos joganyagnak, melyek az MI potenciális kockázatainak kezelésére irányulnak, semelyik ma alkalmazandó jogszabály nem nyújt teljeskörű védelmet a GenMI modellek azon hosszútávú fenyegetettségével szemben, ami a szintetikus tartalmak és az ember által létrehozott autentikus tartalmak megkülönböztetésmentes összeolvadásából származhat.

1.3 Módszertan, a kutatás hasznosítása

A dolgozat elkészítésének alapjául a témához kapcsolódó jelentős hazai és nemzetközi szakirodalom (tanulmányok, monográfiák, kommentárok, tudományos publikációk) áttekintése és feldolgozása szolgált. A kutatási módszertan alapját a releváns Európai Uniós jogszabályok, Tanácsi következtetések és javaslatok, intézményi irányelvek és ajánlások elemzése, valamint különböző esettanulmányok, szakpolitikai jelentések, hatástanulmányok, kockázat- és megfelelésértékelések értelmezése és összefoglalása adta.

A választott téma több tudományterületet (állam- és jogtudományok, politológia, algoritmusetika és szociálpszichológia) érint egyszerre. A dolgozat interdiszciplináris jellege lehetőséget biztosít az MI/GenMI technológiákkal kapcsolatos jelenségek különböző nézőpontokból történő átfogó vizsgálatára. Relevanciáját és kitűzött célját abból az

elgondolásból meríti, hogy az e technológiák széleskörű elterjedéséhez köthető rövid és hosszútávú társadalmpolitikai hatások feltérképezése, illetve az azokkal kapcsolatos szabályozási kihívások részletes tanulmányozása hasznos és szükséges információval szolgálhat majd az illetékes döntéshozók számára. Minden e területen végzett kutatás, megkezdett diskurzus értékes iránymutatásokat és kereteket nyújt a hatékony szabályozáshoz, biztosítva, hogy az MI jövőbeni fejlesztése, forgalmazása és felhasználása etikus, felelős és a demokratikus állam berendezkedések és egyetemes emberi jogok társadalmi értékeivel összhangban lévő legyen. Ez kiváltképp igaz azon GenMI modellekre, melyek relatíve új területnek képviselnek az MI technológiacsalád területén és amelyekkel kapcsolatosan mind a gyakorlati, mind a jogalkotási tapasztalatok korlátozott mennyiségben állnak csak rendelkezésre.

1.4 A dolgozat központi kérdései, hipotézisei

A kockázatalapú megközelítésen alapuló Uniós szabályozás érdeme, hogy helyesen mérte fel a nagy és elfogadhatatlan kockázatúként kategorizált MI alkalmazások jövőbeli potenciális veszélyeit. Az Uniós jogalkotó szervek az elmúlt két évben számos alkalommal előremutató módon voltak képesek pontosítani a jogszabálysöveget e gyakorlatok meghatározásával és besorolásával kapcsolatban. Például, hallgatva az Európai Adatvédelmi Testület és az európai adatvédelmi biztos 5/2021. sz. közös véleményére, helyesen ismerte fel a Tanács, hogy az általános társadalmi pontozás gyakorlatának tilalmát a közszektorról a magánszereplőkre is ki kell terjeszteni. Szintén pozitív fejleményént értékelhető, hogy a végleges szöveg tartalmazza a Régiók Európai Bizottságának azon korábbi javaslatát, miszerint a magatartást torzító MI gyakorlatok tilalmát az életkor, és a szellemi és testi fogyatékoságon túl azon személyek körére is ki kell terjeszteni, akiknek sebezhetősége gazdasági kiszolgáltatottságukból fakad.

Azonban a kockázatalapú szabályozási paradigma - ami technológia specifikus felhasználási eseteinek szabályozására összpontosított - nem tekinthető hatékony megközelítésnek az alapmodellekre jellemző többcélú és dinamikus felhasználási módozatok szabályozására. Az EU rendelet-tervezete kezdetben nem tartalmazott kötelezettségeket a GenMI típusú alapmodellekre vonatkozóan azok újdonsága és a technológia relatív ismeretlensége miatt.

A technológiai változásra a Tanács reagált elsőként, így a 2022 december 6-án közzétett közös álláspontjában egy teljesen új MI kategóriának a bevezetését javasolta általános célú MI néven. Ezt követően a Parlament június 14-én elfogadott kompromisszumos szövegében javaslatot tett

két új típus- az alapmodell és a GenMI – fogalmainak megalkotására. Az elsőt úgy definiálta, mint olyan MI modellt, *amelyet sokoldalúságra terveztek, különféle feladatok elvégzésére képesek és széles körű adatforrásokon képeztek*, az utóbbit pedig olyan MI rendszerként határozta meg, *amely összetett tartalmakat, például videót, hangot és kódot állítanak elő, különböző autonómia szinteken*. A 2023 őszén zajló háromoldalú egyeztetések során a legnagyobb vita az alapmodellek és az általános célú MI-rendszerekkel kapcsolatosan alakult ki. Bár Franciaország, Olaszország és Németország a túlzott szabályozás innovációkat bénító hatását hangsúlyozta, az Unió jogalkotói a szigorúbb szabályozást szorgalmazták, és végül egy kétszintű szabályozási keretet elfogadása mellett döntöttek. Ennek értelmében a törvényszöveg két különböző típusú általános célú MI-rendszert határozott meg, amelyekre eltérő szabályok és jogi kötelezettségek vonatkoznak. Az első kategóriát a normál alapmodellek alkotják, melyek alacsonyabb és kevésbé kiterjedt kockázatot hordoznak, míg a második kategóriát az un. *nagy hatású rendszerszintű kockázatot jelentő alapmodellek képezik*. Az ilyen modellek szolgáltatóira szigorúbb kötelezettségek vonatkoznak majd, amelyek magukban foglalják a modellek értékelését, a rendszerszintű kockázatok felmérését és mérséklését, az ilyen kockázatokról szóló jelentéstételt a Bizottságnak, támadhatósági tesztek végrehajtását, valamint a magas szintű kiberbiztonsági védelem biztosítását. A dolgozat első része többek között ezen jogszabályi kötelezettségek jövőbeli hatékonyságát igyekszik felmérni.

A disszertáció második fele arra keresi a választ, hogy miként kezelhető a GenMI modellek felhasználásával létrehozott példátlan mennyiségű szintetikus tartalom közeljövőben várható gyors és tömeges elterjedése. A dolgozat tézise, hogy az elmúlt három évtizedben a kollektív valóságérzékelést a párhuzamos nyilvánosságok kialakulása és az információs környezet fokozatos perszonalizációja számottevően torzította. A napjainkra jellemző politikai csoportpolarizáció, dezinformációs kitettség, és a hagyományos intézményrendszerekbe vetett bizalom általános csökkenése mind tünetként értelmezhetőek e folyamatban. Azonban a szöveges és audiovizuális tartalmak előállítására alkalmas szabadon elérhető GenMI modellek elterjedése hamar e trend kezelhetetlen eszkalálódásához vezethet.

Az EU a szintetikus tartalmak szabályozása tekintetében az állhírek és politikai deepfake tartalmak, a személyre szabott dezinformációs kampányok, valamint a manipulatív gyakorlatok visszaszorítására összpontosít. Bár ezek fontos lépések az MI-hez köthető ártó gyakorlatok korlátozásában, mégsem nyújtanak elegendő védelmet a mesterségesen előállított tartalmak elterjedését illetően. Az MI-rendelet 50. cikk (4) bekezdése például jól érzékelteti, hogy az

esetek egy jelentős részében az állampolgárok számára nem garantált az a jog, hogy tájékoztatassák őket arról, ha szintetikus tartalmakat fogyasztanak.

Jelen dolgozat feltételezése szerint a szintetikus tartalmak elterjedésének legnagyobb hosszútávú kockázata nem a hibrid hadviselés felerősödése, a választások befolyásolása, tömeges manipulálás és politikai polarizáció fokozódása (ezek rövidtávú veszélyek) hanem az *episztemológiai összeomlás*. Annak a veszélye, hogy a társadalom tagjai végérvényesen elveszítik a bizalmukat az általuk olvasott hírek, információk, tudományos alapvetések és korábban általános konszenzust élvező történelmi események hitelességében. Márpedig a tényekkel kapcsolatos egyetértés, a demokratikus intézményekbe vetett közbizalom és az azonos referenciapontokkal bíró konstruktív politikai vita képes egyedül biztosítani a társadalmak szükséges alkalmazkodóképességét az előttük álló globális átalakulások sikeres navigálásához.

2 Az MI (rövid) története

2.1 Theseus-tól a ChatGPT-ig

Az MI fejlődéstörténetének gyökerei az 1940-es és 1950-es évekig nyúlnak vissza. Ekkor vették kezdetét azon alapvető elméleti tudományos kutatások - kiváltképp a matematika, a logika és a számítástudományok terén - melyek később megalapozták a ma ismert MI technológiáknak.⁶ A korai számítógépes rendszerek közül kiemelendő a Claude Shannon amerikai mérnök-matematikus által 1950-ben megalkotott *Theseus* mechanikus gépegér, mely rendszer különlegessége abban rejlett, hogy nemcsak eljutott a kísérleti labirintusának kijáratához, de a beépített memóriájának köszönhetően el is tudta raktározni az előzőleg elvégzett műveletet sorozatot - másképpen megfogalmazva „emlékezett” arra az útvonalra, amit a kezdőponttól a célállomásig megtett.⁷

Ugyanezen évben, 1950-ben közölte nagyhatású cikkét "*Computing Machinery and Intelligence*" címmel Alan Turing brit matematikus. A cikkben Turing egy gondolkísérlet erejéig felvázolt egy általa kitalált játékot (Imitation Game, kisebb változtatásokkal ez nevezzük ma Turing-tesztnek), ami arra szolgált, hogy megállapítsa képes-e egy gép olyan *intelligens* módon kommunikálni, amely megkülönböztethetetlen az embertől. A teszt három résztvevőjét - egy számítógépet és két embert (válaszoló / kérdező) - külön szobákba helyeztek. A kérdező tudta, hogy az egyik résztvevő gép, a másik ember, de arról nem volt tudomása, hogy melyik válaszoló melyik. A kérdező – szöveges csatornán keresztül - kérdéseket tett fel az emberi résztvevőnek és a számítógépnek is, hogy megállapítsa, melyik a válaszoló az ember és melyik a számítógép. A gép célja olyan kimenetek generálása volt, amelyek megkülönböztethetetlenek az emberi válaszoktól. Ha a kérdező nem tudja megbízhatóan megkülönböztetni a gépet az

⁶ Példaként – némi elfogultsággal – említhető a magyar származású matematikus és fizikus, Neumann János munkássága, akinek 1945-ben publikált híres modellje – a Von Neumann-architektúra - a modern számítógépek felépítésének máig érvényes elméleti alapjait fektette le. A Neumann-modell alapján egy számítógép öt fő komponensből áll, melyek a: (1) Feldolgozó Egység (CPU); (2) Memória; (3) Bemeneti Egység (Input); (4) Kimeneti Egység (Output); (5) Vezérlő Egység (Control Unit). Lásd bővebben: Von Neumann, J. (1945): First Draft of a Report on the EDVAC. Az eredeti szöveg beszkenelt formában angolul elérhető: https://archive.computerhistory.org/resources/text/Knuth_Don_X4100/PDF_index/k-8-pdf/k-8-u2593-Draft-EDVAC.pdf

⁷ Ennek a képességnek köszönhetően a gépegér már nem véletlenszerűen tájékozódott a labirintusban, hanem a korábban bejárt útvonalra alapozva optimalizálta mozgását, lecsökkentve a labirintusból kivezető út hosszát. Jóllehet ma már kezdetlegesnek tűnhet e rendszer, jelentősége mégsem elhanyagolható, hiszen a Theseus volt az egyik első példája annak, hogy gépek bizonyos szinten megtanulhatják és memóriájukba vészik az elvégzett feladatokat. Bővebben: Klein, Daniel (2018): Mighty mouse. *MIT Technology Review*. Elérhető: <https://www.technologyreview.com/2018/12/19/138508/mighty-mouse> (2018. 12. 18.)

embertől, akkor a gép átment a teszten.⁸ Bár cikk révén népszerűvé váló Turing-tesztet az évek során sok kritika érte, elsősorban a kommunikációs képesség magas szintű utánzása és az absztrakt emberi intelligencia párhuzamba állítása okán, mégis az imitációs játék gondolat kísérlete és maga Turing is tartós hatást gyakorolt arra miként határozták meg sokáig az intelligens rendszerek kritériumait.

Az MI kifejezést elsőként John McCarthy amerikai matematikus, a Stanford egyetem professzora használta szélesebb közönség előtt az 1956-ban megtarott Hampshire-i Dartmouth Egyetem nyári workshopján. Az egy évvel korábban publikált kutatási projektjavaslatukban McCarthy és kollégái olyan jövőbeni irányvonalakat fogalmaztak meg a gépi intelligencia területén végzett kutatásokban, melyek főként annak problémamegoldási képességeire összpontosítottak (automatizáció, önfejlesztés képessége, kreativitás, természetes nyelv feldolgozása, elvont fogalmak értelmezése).⁹ Többek között a Dartmouthi Egyetem kutatóihoz köthető diszciplináris keretrendszernek köszönhető, hogy a fogalom olyan gépekkel összefüggésen vált széles körben használatossá, melyek az emberi intelligenciát igénylő feladatokat és problémák elvégzésére válhat alkalmassá.

Annak ellenére, hogy az MI területén végzett kutatások jelentős áttöréseket hoztak¹⁰, a korai ígérek túlságosan nagyratörőnek és optimistának bizonyultak, így fokozatos csalódottság és szkepticizmus alakult ki a befektetők körében, ami az 1970-es években a technológiai kutatások, innovációkkal kapcsolatos általános érdeklődés csökkenéséhez, ún. "MI télhez" vezetett.¹¹ Nem telt el sok idő azonban, és 1980-as évek második felétől kezdődően újra nőni

⁸ Lásd bővebben: Turing, A. M. (1950): Computing Machinery and Intelligence. *Mind*, 59(236), 433–460.o. Az eredeti cikk szövege magyarul elérhető Tarján Rezsőné fordításában: <https://www.geier.hu/GOEDEL/Turing/Turing.html>

⁹ A projekt során a következő hét alapvető kutatási területet határozták meg: (1) Automatizált számítógépek, melyek emberi beavatkozás nélkül is képesek lehetnek különböző feladatok elvégzésére; (2) Olyan modellek vizsgálata, mely lehetővé teszi a számítógépek számára az emberi természetes nyelvek megértését és használatát; (3) Olyan az emberi agy szerkezetét és működését utánzó neurális hálózatok tanulmányozása, amelyek képesek a tanulásra és a mintafelismerésre; (4) Az algoritmusok, számítási modellek és számításelméleti problémák vizsgálata a számítógépek elméleti képességeinek és korlátjainak feltérképezésére; (5) Kutatások a gépi önfejlesztés és öntanulás területein (gépi tanulás) (6) A feldolgozott adatokból történő absztrakcióképzés gépi módszereinek leírására (7) Kutatások a kreativitás gépi szimulálásával kapcsolatban, főként arra vonatkozóan, hogy miként generálhatnak a számítógépek új ötleteket vagy megoldásokat. A kutatási projekt résztvevői: John McCarthy, Marvin Minsky, Allen Newell, Arthur Samuel, Herbert Simon. A kutatási javaslatról bővebben: McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1955): *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*. Az eredeti szöveg angol nyelven elérhető: <https://www.formal.stanford.edu/jmc/history/dartmouth/dartmouth.html> (1996. 04. 03.)

¹⁰ Rövid ismertetés a tudásalapú és szakértői rendszerekről XXX

¹¹ Az első meghatározó MI tél a 1970-es évek közepén zajlott, amit több kisebb visszaesés is követett az 1980-as és a 2000-es években. (A legkorábbi – igaz kevésbé meghatározó – MI telet az 1950-es évekre datálják, amikor is egy viszonylag egyszerű, orosz tudományos szövegeket angolra fordító program kezdeti sikerén felbuzdulva, sokan az intelligens beszélő rendszerek korai eljövételét vizionálták). Ezekben az időszakokban a

kezdt a lelkesedés a technológiacsalád iránt, majd következő évtized elejére a gépi tanulás fejlődésével és a neurális hálózatokkal kapcsolatos kutatások megindulásával beköszöntött az „MI tavasz”.¹²

Az 1990-es években az internet elterjedése és a szabadon hozzáférhető digitális adatmennyiség exponenciális növekedése további új lehetőségeket nyitottak a kutatások területén. Ebben az időszakban született meg a IBM által fejlesztett Deep Blue sakkszámítógép is, amely 1997 májusában egy hatjátszmás páros mérkőzésen 3,5-2,5 arányban győzte le Garry Kasparovot, az akkori sakkvilágbajnokot.¹³ Az emlékezetes meccs nemcsak a sakk, hanem az ember-gép interakció történetében is jelentős pillanat volt; ez tekinthető az első olyan eseménynek, melynek köszönhetően a társadalom számottevő része felfigyelt az MI-ben rejlő potenciális képességekre. Az IBM által tervezett Deep Blue működése még nem gépi tanuláson vagy neurális hálózatokon alapult, pusztán hatalmas számítási kapacitására és az előre beprogramozott sakkstratégiákra támaszkodott. Erősségét innovatív (alfa-béta) keresőalgoritmusra és célhardver adta, melynek köszönhetően másodpercenként megközelítőleg 200 millió pozíciót is képes volt áttekinteni, szisztematikusan végig számolva a lehetséges lépéseket.¹⁴

A következő nagyobb visszhangot kiváltó esemény szintén egy IBM által fejlesztett MI-vel, a Watson nevű természetes nyelvfeldolgozó rendszerrel (Natural Language Processing, röviden NLP) volt kapcsolatos, amely 2011-ben két emberi versenyzővel megmérkőzve ért el első helyezést a "Jeopardy!" nevű amerikai kvízműsorban. A gép megértette az emberi nyelven megfogalmazott kérdéseket, és a kiterjedt adatbázisában tárolt információk alapján képes volt rendkívül alacsony hibaaránytal helyes válaszokat generálni azokra. Ez a siker rávilágított, hogy a MI rendszerek immár alkalmasak lehetnek a nyelvfeldolgozás és általános tudás alapú feladatok ellátására is.

korábbiakhoz képest lassabb ütemben haladt a kutatás, több projektet félbehagytak, és számottevően csökkent a tudományterületekre szánt finanszírozás is. Bővebben: Floridi, Luciano (2020): AI and Its New Winter: from Myths to Realities. *Philosophy & Technology*. 33, 1–3.o. <https://doi.org/10.1007/s13347-020-00396-6>

¹² Suleyman, Mustafa, Michael Bhaskar (2023): A következő hullám. Mesterséges intelligencia, technológia, hatalom és a 21. század legnagyobb kihívása. (ford. Farkas Veronika). Budapest, Magnólia kiadó. 74-76.o.

¹³ A Carnegie Mellon Egyetem kutatói a 80-as évek közepén kezdték el fejleszteni Deep Thought (korábban ChipTest) nevű sakktanulmányprogramot, amely az évtized végére, 1989-ben legyőzte a dán Bent Larsen sakknagymestert. E kutatásokba szállt be később az IBM is, és fejlesztette ki saját programját – immár Deep Blue néven. Garry Kasparov az elhíresült mérkőzést megelőző évben 1996-februárjában még képes volt – viszonylag könnyen - 4-2-es arányban legyőzni ezt a programot. Bővebben: Hassabis, D. (2017): Artificial Intelligence: Chess match of the century. *Nature* 544. 413–414.o. <https://doi.org/10.1038/544413a>

¹⁴ Campbell, M., Hoane Jr, A. J., & Hsu, F. H. (2002). Deep Blue. *Artificial Intelligence*, 134(1-2), 57–83.o.

A 2010-es évektől kezdődően az MI fejlődését elsősorban már a neurális hálózatok¹⁵ és mélytanulási módszerek elterjedése, valamint a nagy adathalmazok (Big Data)¹⁶ rendelkezésre állása fűtötte. A mélytanulás - ahogy arról a következő fejezetben szó lesz - lehetővé tette a gépek számára, hogy összetettebb mintákat és kapcsolatokat ismerjenek fel az adatokban, míg a Big Data biztosította a szükséges adatmennyiséget a modellek hatékony betanításához. Ezen technológiák sikeres adaptálásáról tanúskodik a Google DeepMind elsőpró sikere is.

2017 decemberében a Google DeepMind által fejlesztett AlphaZero program azzal vonta magára a tudományos világ figyelmét, hogy egy száz játszmából álló mérkőzésen legyőzte a 2016-os számítógépes sakk-világbajnokot, Stockfish8-at. Az Alphazero mindössze négy óra alatt sajátította el a sakk szabályait, anélkül, hogy a hagyományos értelemben programozták, vagy „tanították” volna a játékra. A Google fejlesztése a saját maga ellen folytatott több millió játszma során gyűjtött adatok (pl.: sikeres - sikertelen lépéskominációk) alapján folyamatosan finomította és fejlesztette saját stratégiáját. A tapasztalatokat beépítve a neurális hálózatába, az

¹⁵ Neurális hálózatok az 1980-as évek vége felé kezdtek megjelenni az akadémiai kutatásokban, és 90-es években kezdtek el komoly figyelmet kapni, majd a 2000-es évek kezdtek elterjedni a gyakorlati alkalmazások a természetes nyelvfeldolgozás terén

¹⁶ A hétköznapi értelemben a Big Data, vagyis nagy adatmennyiség kifejezés egyszerre utal a méretükben és összetettségükben jelentős adatkészletekre, valamint arra technológiai eszköztárra, amely képes feldolgozni azt. Kezdetben a fogalmat - Douglas Laney 2011-es publikációját követően - az adatok három V-je segítségével határozták meg, ami a Volume (mennyiség), Variety (változatosság), és Velocity (sebesség) szavakat jelölte; (1) A hatalmas mennyiség a Big Data talán legfontosabb jellemzője. Manapság az egyének, vállalatok, államigazgatási szervek egytől egyig nagy mennyiségű adatot állítanak elő online/ offline tevékenységeiken keresztül. Az előállított adatok robbanásszerű növekedése nem kis mértékben köszönhető a mobil és egyéb digitális eszközök, az IoT szenzorok, a közösségi média és más platformszolgáltatások elterjedésének is. (2) A Sebesség az adatok keletkezésének és áramlásának a gyorsaságára utal. A valós idejű adatforrások, mint az érzékelők, streamingszolgáltatások, mobil eszközök és online tranzakciós rendszerek folyamatosan ontják az új adatokat bármilyen interakció során. (3) A Változatosság az adatok típusainak és formátumainak sokféleségére utal. A Big Data esetében a legkülönbözőbb forrásokból érkező adatok lehetnek strukturáltak (adatbázis táblázatok), félig strukturáltak (XML dokumentumok), vagy teljesen strukturálatlanok (szabad formátumú szövegek, mint például a felhasználói hozzászólások egy komment szekcióban). Bár ez a változatosság bonyolulttá teszi az adatkezelést és az adatintegrációs és adatelemzési feladatokat, mégis nagy potenciált kínál a legkülönbözőbb területekről begyűjtött információk hasznosítására. Lásd: Gartner (2012): Gartner Research. The Importance of 'Big Data': A Definition. Elérhető: <https://www.gartner.com/en/documents/2057415> (2020. 07. 12). De a Big Data koncepció kiegészülni látszik további 2 V-vel is, melyek a Veracity (Hitelesség) és Value (Érték) szavakat takarják. Az első V arra utal, hogy a nagy adathalmazok esetén gyakran nehéz megbizonyosodni azok megbízhatóságáról és pontosságáról, mivel azok sokszor eltérő forrásokból és változó minőségben érkeznek. Az Érték pedig azt a felismerést tükrözi, hogy az adatokban rejlő információk - megfelelő adatelemző eljárásokkal - valódi gazdasági vagy társadalmi hasznot rejtenek. (Tovább bonyolítandó, de ehelyütt részletesen nem tárgyalva Eileen McNulty 2014-ben már V7ről is értekezett, amely a Variability és Visualisation jellemzőkkel történő kibővítéssel jönne létre. McNulty, Eileen (2014): Understanding Big Data: The Seven V's. DataConomy. Elérhető: <https://dataconomy.com/2014/05/22/seven-vs-big-data> (2014. 05. 22.) Témáról bővebben: Viktor Mayer-Schönberger - Kenneth Cukier (2014): Big data. Forradalmi módszer, amely megváltoztatja munkánkat, gondolkodásunkat és egész életünket. Ford. Dankó Zsolt. Budapest, HVG Könyvek; Szőke Gergely László (2018): Big Data and Algorithms in the Public Sector and Their Impact on the Transparency of Decision-Making. Central and Eastern European eDem and eGov Days. 301-306.o.; Zódi Zsolt: Jog és jogtudomány a Big Data korában. Állam- és Jogtudomány 2017/1.

AlphaZero képessé vált a játék mélyebb megértésére és új, kreatív stratégiák kifejlesztésére.¹⁷ Ez a fajta *megerősítéses tanulás* lehetővé tette az AlphaZero számára, hogy saját maga fedezze fel a játék mélységeit és új, kreatív lépéseket alkalmazzon. Az AlphaZero megközelítése radikális eltérést jelentett a hagyományos, adatbázis-alapú MI programoktól, mint amilyen a Deep Blue volt.¹⁸ Míg a Deep Blue a hatalmas adatbázisokra és előre betáplált stratégiákra támaszkodott, az AlphaZero teljesen önállóan, a mély tanulás és a neurális hálózatok segítségével fejlesztette ki saját játékstratégiáit anélkül, hogy előre beprogramozott vagy emberi sakkozók által kikísérletezett lépéssorozatokot használt volna.

A 2022-es év újabb mérföldkönek tekinthető az MI fejlődésében, különösen a generatív MI (GenMI) modellek területén. Ebben az évben vált széles körben elérhetővé és ingyenesen hozzáférhetővé a nagyközönség számára a ChatGPT Nagy Nyelvi Modell (Large Language Model, LLM), amely máig e modellcsalád talán legismertebb képviselője. Ezek alapjául a későbbiekben még tárgyalandó a GAN (Generative Adversarial Networks) és a transzformer-alapú mélytanulás modellek szolgálnak. A szoftver megjelenése óta alig telt el több mint másfél év, de a lelkesedés az MI rendszerek új technológiája iránt töretlen. Kutatók, fejlesztők és felhasználók egyaránt keresik a módját, hogyan lehet a GenMI modelleket különböző területeken alkalmazni, legyen szó szövegírásról, képgenerálásról, kódolásról vagy éppen adatelemzésről. A technológiai területen gyakran hivatkozott Gartner-féle hype-ciklust¹⁹

¹⁷ A Deepmind fejlesztései nem csak a sakk, hanem a go és a sógi játékokban is képesek voltak új stratégiák kidolgozására és a világ legjobb játékosainak legyőzésére. 2016. március 9-15. között, a szintén DeepMind által fejlesztett AlphaGo megmérkőzött Lee Sedol-lal, a világ egyik legjobb go játékosával, akit 4:1 arányban legyőzött. Ez az esemény azért volt figyelemreméltó, mert a go játék rendkívül összetett, a lehetséges lépések és állások száma emberi tudattal szinte felfoghatatlan (2×10^{170}), amely hatalmas kihívás a számítástechnikai kapacitások és keresőalgoritmusok szempontjából. Nem is csoda, hogy az áttörést végül a Google mélytanulás módszerei hozták el. A mérkőzés után Lee Sedol azt nyilatkozta: XXX

¹⁸ A Deep Blue programot még úgy tanították, hogy millió létező sakklépést és játékstratégiát tápláltak bele. Ennek köszönhetően a lehetséges (ismert és valós mérkőzéseken alkalmazott) lépések széles körét ismerte, mielőtt szembe került volna egy emberi ellenféllel. Ez a megközelítés bár fejlett adatelemzést és jelentős számítástechnikai teljesítményt igényelt, mindig kizárólag csak arra a játéktípusra lesz alkalmazható amire előzetesen betanították. Lásd: Campbell, M., Hoane Jr, A. J., & Hsu, F. H. (2002). Deep Blue. *Artificial Intelligence*, 134(1-2), 59–60.

¹⁹ A Gartner féle hype-ciklus (Gartner Hype Cycle) szerint a technológiai innovációk fejlődésében és fokozatos elterjedésében öt szakasz különböztethető meg: (1) Technológiai áttörés (Technology Trigger): Egy adott technológia korai sikertörténetei okán az iránta tanúsított érdeklődés aránytalanul nagyra duzzad, a közbeszédet is uralja a téma. A média által nagy nyilvánosságot kap a terület, bár gyakran még sem működő prototípusok, sem létező gazdasági modell nem épült fel a technológia körül. (2): A felfokozott várakozások csúcsa (Peak of Inflated Expectations): Lényegesen nő az új technológiát ismerők és amatőr szinten használók száma. Egymást érik a fantasztikus áttörésekről szóló gyakran szenzációhajhász hírek. A technológia iránti lelkesedés a csúcspontjára ér, sokszor irreális elvárásokkal és túlzott optimizmussal. (3) A csalódás görbéje (Trough of Disillusionment): A felmerülő problémák növekednek, a forradalmi áttörések száma csökken. Amikor a technológia nem tudja teljesíteni a magas elvárásokat, csalódás következik be és az érdeklődés csökken. Kutatások és a becsatározott tőke is lényegesen apad. A technológia akár zsákutcába is fordulhat. (4): A megvilágosodás emelkedője (Slope of Enlightenment): A technológia valós lehetőségeinek és korlátjainak fokozatos felismerése következik. A korai tapasztalatok születte módosítások és a használat közben felgyülemlett tudás révén lassan kikristályosodik a technológia használhatóságának köre és módja. Bár a befektetők nagy része még óvatos, kezd kialakulni egy stabil,

figyelembe véve-, mely az innovációk és az azzal kapcsolatos társadalmi attitűdök jól megfigyelhető trendjét osztja öt szakaszra - joggal feltételezhető, hogy a GenMI modellekkel kapcsolatban a társadalmi reakciók a ciklus második fázisánál tartanak. Erre az időszakra jellemző a feltűnő és médiatértben is hangos, átütő sikerek miatti hatalmas felhasználói és médiafigyelemmel övezett, folyamatosan növekvő és gyakran alaptalan lelkesedés. Kérdés azonban, hogy mikor éri el – ha egyáltalán eléri – a harmadik szakaszt, amely jellemzően az új technológiával szembeni fokozatos kiábrándulással és apátiával jár. Mivel a GenMI modellek minden bizonnyal szerves részét fogják képezni a technológiacsatládnak, e téma a disszertáció második részében még részletesebben is kifejtésre kerül.

2.2 Az MI-ről általában

Nem várt nehézségekbe ütközhet az, aki az MI fogalmának átfogó, technológiasemleges meghatározására vállalkozik. Bár az Európai Unió többéves jogalkotási folyamata sikeres erőfeszítésnek tekinthető egy időtálló MI definíció megalkotására – erről a későbbiekben részletesen is szó lesz - valójában mindmáig nincsen egyetlen általánosan elfogadott, egységes meghatározás e rendszereket illetően. Vagy másképpen fogalmazva: számtalan széles körben elfogadott értelmezés létezik párhuzamosan.

Stuart Russell és Peter Norvig *Artificial Intelligence: A Modern Approach* című 1995-ben megjelent könyvükben a következő négy szempont szerint csoportosította az MI-hez köthető kutatások és fejlesztések irányát (1) emberi gondolkodás, (2) racionális gondolkodás, (3) emberi viselkedés, (4) racionális viselkedés. Minden kategória azt vizsgálja, hogy hogyan értelmezhető az intelligens rendszer fogalma, és hogy miként valósíthatók meg a gyakorlatban ezek az elképzelések.²⁰

Az első felfogásban az emberi gondolkodás utánzásán van a hangsúly, ahol olyan gépek létrehozása a cél, melyek az emberi gondolkodási folyamatokat (például: tanulás, emlékezés, problémafelismerés, problémamegoldás) és az emberi megismerés jellemzőit (pl. annak

megbízható felhasználói réteg. (5): A produktivitás fennsíkja (Plateau of Productivity): A technológia eléri azt a szintet, ahol stabilan és hatékonyan használható, és valós értéket képes létrehozni. A termék a hétköznapiak részévé válik, használata elterjed, megtalálja helyét a piacon. Lásd bővebben: Fenn, Jackie; Raskino, Mark (2008): *Mastering the Hype Cycle: How to Choose the Right Innovation at the Right Time*. Massachusetts, Harvard Business Publishing.

²⁰ A négy kategória angolul: (1) Thinking Humanly (2) Thinking Rationally (3) Acting Humanly (4) Acting Rationally. Lásd bővebben: Russell, Stuart J. – Norvig, Peter (2010): *Artificial Intelligence: A Modern Approach*. Third Edition. Essex, Pearson. 2-16.o.

megértése, hogy mások meggyőződése, vágyai és szándékai hogyan befolyásolják döntéseiket) képesek imitálni. A második kategóriában a racionális gondolkodáson van a hangsúly melyben a gépek létrehozásánál a logikai érvelés tökéletesítésére összpontosítanak, olyan gépek fejlesztésére, amelyek képesek racionális döntéseket hozni és logikailag megalapozott következtetéseket levonni. A harmadik típusnál a cél, hogy olyan gépeket hozzanak létre, amelyek képesek az emberi viselkedés utánzására. Például, ha egy gép képes beszélgetni egy emberrel anélkül, hogy az ember rájönne, hogy géppel beszél – amire a Turing-teszt is irányul - akkor intelligens rendszerrel állhatunk szemben. A negyedik kategória esetében a hangsúly a racionális, cél-orientált viselkedés megvalósítására összpontosul. E felfogásban a cél, hogy a gépek képesek legyenek az optimális döntések meghozatalára a rendelkezésre álló információk és célkitűzések alapján.

Ahogy e kategorizálásból kitűnik, egy könnyen érthető és kellően általános megfogalmazásban az MI azon technológiai rendszereket összeségét jelöli, amelyek képesek utánozni az olyan emberi kognitív funkciókat, mint az érzékelés, logikai következtetés, problémamegoldás és a tanulás. Efféle kézenfekvő értelmezést sugall az MIT fizikaprofesszora, Max Tegmark is, aki szerint legáltalánosabb felfogásban az MI-re, mint az emberi intelligenciát imitálni képes rendszerre érdemes tekinteni. Tegmark szerint az intelligencia nem más, mint az *összetett célok elérésének képessége*.²¹ Tartalmát tekintve hasonló, gyakran idézett definíciót alkotott az amerikai számítógép-tudós Nils Nilsson is, aki úgy véli *"az MI olyan tevékenység, amely arra irányul, hogy a gépeket intelligenssé tegye. Az intelligencia pedig olyan tulajdonság, amely lehetővé teszi egy adott entitás számára, hogy megfelelően és előre látóan működjön környezetében"*.²² Ahhoz kötődően, hogy e célokat, milyen mértékben képes elérni egy adott MI, általánosságba véve két fő kategóriát szokás elkülöníteni: a szűk (vagy gyenge) valamint az általános (vagy erős) MI.

2.3 Szűk és Általános MI

A szűk MI (narrow AI) olyan szoftvereket takar, melyek egy adott probléma megoldására vagy egy konkrét feladat végrehajtására specializálódtak. Ezek a rendszerek egy előre meghatározott feladatkörben működnek, és működésüket egy dedikált szabályrendszer határozza meg, amely

²¹ Tegmark, Max (2017): Élet 3.0. Embernek lenni a mesterséges intelligencia korában (Ford.:Weisz Böbe, Garai Attila). HVG Könyvek, Budapest 69.o.

²² „Artificial intelligence is that activity devoted to making machines intelligent, and intelligence is that quality that enables an entity to function appropriately and with foresight in its environment). Lásd: Nils J. Nilsson (2009): The Quest for Artificial Intelligence. Cambridge University Press. 13.o.

informatikai nyelvre van lefordítva. Ezen a behatárolt területen a szűk MI gyakran képes az emberi teljesítményt felülmúlni, ugyanakkor ez a kiemelkedő teljesítmény csak az adott feladat mechanikus végrehajtására korlátozódik, mivel nem rendelkezik komplex problémamegoldó képességgel vagy általános intelligenciával. A szűk MI erőssége egyben a gyengesége is: a specializált feladatra való optimalizálás megakadályozza, hogy ezek a rendszerek az emberhez hasonló, széles körű problémamegoldó képességet érjenek el.²³ Hiányzik belőlük az általános intelligencia, a tanulási képesség és az alkalmazkodóképesség, amelyek az emberi gondolkodás sajátosságai, így a szűk MI rendszerek, bár kiemelkedően teljesítenek egy adott területen, nem képesek az emberhez hasonló módon kezelni az új, váratlan helyzeteket vagy a kontextus változásait. Szűk MI-k például azon rendszerek, melyeket adatelemzésre, kép-és mintafelismerésre vagy szövegértelmezésre használnak

Ezzel szemben általános MI (General AI) egy olyan fejlett intelligencia, amely képes széles körű feladatok elvégzésére, önálló tanulásra és alkalmazkodásra. Ez az intelligencia képes új, ismeretlen problémákat megoldani és különböző területeken kiemelkedően teljesíteni, miközben az emberi gondolkodás sajátosságait tükrözi. Képesek felismerni és megérteni összetett mintázatokat, észrevenni a kontextus változásait, és új problémákat kreatív módon megoldani. Ezen képességek révén az általános MI nem korlátozódik egyetlen feladat mechanikus végrehajtására, hanem képes különböző területeken is kiemelkedően teljesíteni. Ezen kívül fejlett problémamegoldó képességekkel rendelkezik, amelyek révén új, ismeretlen helyzetekben is képes helytállni. Az ilyen rendszerek képesek összetett döntéseket hozni, hosszú távú következményekkel számolni és stratégiai gondolkodást alkalmazni. A fejlett általános MI még a jövő technológiája, de kétségtelen, hogy a következő alfejezetben tárgyalandó gépi tanulás eljárások jelentős előre lépéseket jelentettek fejlesztéséhez vezető úton.

E kettő kategóriát gyakran ki szokták egészíteni egy harmadik – fiktív – kategóriával is, melyet Mesterséges Szuperintelligenciának hívnak. Ez a típusú MI minden szempontból felülmúlja az emberi intelligenciát, a felhalmozott tudástól kezdve a kreativitáson át, a problémamegoldó képességekig. Nick Bostrom, az Oxfordi Egyetem filozófusa a Szuperintelligencia című könyvében az fejlődésének három fő szintjét írta le (1): Gyors Szuperintelligencia: olyan rendszer, ami mindenre képes amire az ember, csak sokkal gyorsabban (2): Kollektív szuperintelligencia: nagyszámú kisebb intelligenciából felépített rendszer, amelynek összesített

²³ Klein Tamás (2021): Robotjog vagy emberjog? Az emberközpontú mesterséges intelligencia szabályozásának kiindulási pontjai. In Török Bernát és Zódi Zsolt (szerk) A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 132.o.

teljesítménye számos általános területen jóval túllépi a mai rendszerek teljesítményét (3):Minőségi szuperintelligencia: olyan rendszer, amely legalább olyan gyors, mint egy emberi elme, de képességeiben messze meghaladja az emberi intelligencia határait.²⁴

2.4 Az öntanuló és önfejlesztő MI

Az MI gyűjtőfogalom számos különböző technológiát, eljárást és módszert foglal magában. Jelen dolgozatnak nem kitűzött - és területi korlátjai okán, nem is megvalósítható - célja ezek részletes ismertetése. Azonban a modern értelemben vett MI pontosabb megértése érdekében mégis indokolt egy a technológiacsaládhoz szorosan köthető alapfogalom - a gépi tanulás - kifejtése és főbb típusainak rövid bemutatása.

A gépi tanulás az MI egy részterülete, amely kifejezetten az adatokból és tapasztalatokból való tanulásra koncentrál. A gépi tanulás térhódítása az 1990-es években kezdődött, és rövidesen forradalmasította az MI képességeit azáltal, hogy lehetővé tette a rendszerek számára, hogy önállóan (iteratív módon) tanuljanak és fejlődjenek, anélkül, hogy külön programozásra lenne szükségük minden egyes feladat vagy probléma megoldásához. A gépi tanulás elterjedését megelőzően a hagyományos számítógépes alkalmazások működése mereven kötött volt a programozók által előre megírt utasításokhoz. Ezek az instrukciók (programkódok) lépésről lépésre meghatározták, hogy a gépnek pontosan mit kell tennie egy adott feladat végrehajtásához. Bármennyire is kifinomultak voltak ezek a programok, csak arra voltak képesek, amire explicit módon betanította őket a forráskód.²⁵ Ezzel szemben a modern MI rendszerek olyan rugalmasabb megközelítést alkalmaznak a gépi tanulás révén, ahol nincs szükség részletes és folyamatos (újra)programozásra. Ehelyett hatalmas mennyiségű nyers adatot táplálnak be a rendszerbe a legkülönbözőbb forrásokból. Az algoritmusok²⁶ azután

²⁴Bár ez a kategória egyelőre csak elméleti síkon létezik, mégis egyre több MI kutató, társadalomtudós és filozófus fejezi ki aggodalmát e magasan fejlett technológiacsaláddal kapcsolatban. A témáról lásd bővebben Bostrom, Nick (2015): Szuperintelligencia. Utak, veszélyek, stratégiák. ford. Hidy Mátyás. Budapest, Ad Astra. 89-96.o.

²⁵ Például a szakértői rendszerek. XX

²⁶ Az algoritmus az MI alapját képező, lépésről lépésre végrehajtandó matematikai utasítások sorozata. Ezek az utasítások határozzák meg, hogy a rendszer hogyan dolgozza fel az adatokat és hoz döntéseket. Az algoritmusok rendelkeznek bemeneti (input) és kimeneti (output) adatokkal; a bemeneti adatok felhasználásra kerülnek valamilyen művelet során, és az eredmény az algoritmus végén kimeneti adatként jelenik meg. Az algoritmusok terjedhetnek egyszerű szabályalapú rendszerektől az összetett prediktív modellekig. Könnyen érthető példaként szolgálhatnak az online áruházak által előszeretettel alkalmazott személyre szabott ajánlórendszerek. Egy ilyen rendszer az alábbi lépésekben (utasítási sorozatokban) működik: Először elemezni kezdi a felhasználó korábbi vásárlásait és böngészési adatait. Ezen információk alapján meghatározza, milyen termékek iránt érdeklődhet leginkább a felhasználó. Ezután keres a termékdatabázisban olyan cikkeket, amelyek megfelelnek a felhasználó

képesek felismerni és elemezni az adatokban rejlő struktúrákat, mintázatokat és összefüggéseket. Ez a tanulási folyamat teszi lehetővé, hogy az MI kikövetkeztesse a megfelelő működési módot, és új, ismeretlen helyzetekben is helyesen alkalmazhassa megszerzett tudását.

A gépi tanulás talán legismertebb ága a mélytanulás (deep learning). Ez a módszer az emberi agy működését utánzó, többrétegű neurális hálózatokon alapul, ahol minden réteg külön-külön dolgozza fel a bemeneti adatokat, és ezeket az információkat továbbítja a következő rétegek felé.²⁷ A 'mély' kifejezés a sok rejtett rétegre utal a hálózatokban. A neuronok közötti súlyozott kapcsolatok révén a hálózat képes összetett mintázatokat tanulni és előre jelezni anélkül, hogy explicit, feladat-specifikus programozási szabályokra lenne szükség.²⁸ A mélytanulási hálózatok többféle formában léteznek, ismert típusai például az egyrétegű, az előrecsatolt, az önszervező, a konvolúciós (CNN), és a rekurrens (RNN) neurális hálózatok.²⁹

A technológiára a 2010-es évek elejétől irányult kiemelt figyelem többek között annak köszönhetően, hogy Alex Krizhevsky, Ilya Sutskever és Geoffrey Hinton 2012-ben bemutatta mély konvolúciós neurális hálózatát, (deep convolutional neural networks, CNN) amely példátlan mértékben, 28%-ról 15,3%-ra csökkentette a hibaarányt az MI képfelismerési feladatokban az addigi legjobb eredményhez képest.³⁰ Az elkövetkező években neurális

ízlésének és igényeinek. Végül a rendszer rangsorolja és megjeleníti a felhasználó számára leginkább releváns termékeket. Összefoglalva tehát az algoritmus megengedett lépésekből, tevékenységekből álló véges utasítássorozat, amelyet követve elérhetjük a meghatározott, kívánt célt. A témáról részletesebb leírást nyújt a Debreceni Egyetem Informatikai Karának oktatási tananyaga: Varga Imre (2020): Algoritmusok és a programozás alapjai. Elérhető: <https://irh.inf.unideb.hu/~vargai/APA/index.html>

²⁷ Aneurális hálózatok bemeneti, rejtett és kimeneti rétegekből állnak: a bemeneti réteg feladata az adatok fogadása és továbbítása a hálózat belső rétegei felé. A rejtett rétegekben történik az adatok tényleges feldolgozása, átalakítása. Itt a beérkező információk különböző matematikai műveletek segítségével új reprezentációkká formálódnak, miközben a hálózat kiemeli a lényeges jellemzőket és mintázatokat. Végül a kimeneti réteg szolgáltatja a feldolgozás végeredményét a hálózat aktuális céljának megfelelő formában. A neurális hálózat minden rétegében mesterséges neuronok találhatók, melyek a súlyokkal és aktivációs függvényekkel vezérlik az adatok áramlását és feldolgozását. Az aktivációs függvény határozza meg, hogy egy neuron aktiválódjon-e, és ha igen, milyen mértékben továbbítsa a jeleket a következő réteg felé. A hálózat tanulása során a súlyok és a bias értékek folyamatosan módosulnak annak érdekében, hogy a bemeneti adatokból minél pontosabban előálljon a kívánt kimenet. Például, ha egy képelemző MI feladata, hogy rubik-kockákat ismerjen fel, a tanítás során kép pixeladatait a bemeneti réteg veszi fel, majd ezeket az adatokat a rejtett rétegek több lépcsőben dolgozzák fel, ahol a neuronok aktivációs függvények segítségével a tárgy azonosítására szolgáló jellemzőket emelik ki (forma, soronként eltérő színek stb.). A folyamat végeztével a kimeneti réteg dönti el, hogy látható-e az adott képen rubik-kocka. A tanulási folyamat során az ún. backpropagation algoritmus folyamatosan finomítja a hálózat súlyait a pontosabb felismerés érdekében. Lásd részletesebben: Schmidhuber, J. (2015): Deep learning in neural networks: An overview. *Neural Networks*, 61, 85-117.o.

²⁸ LeCun, Yann, Yoshua Bengio és Geoffrey Hinton (2015): Deep learning. *Nature* 521 (75539), 436-444.o.

²⁹ Gyires-Tóth Bálint (2020): A mélytanulás múltja, jelene és jövője. *Híradástechnika*. 75/1., 23-29.o.

³⁰ A hálózat megközelítőleg 60 millió paraméterrel és 650 000 neuronnal rendelkezik, öt konvolúciós rétegből áll, melyeket pooling (összegző) rétegek, majd három teljesen összekapcsolt réteg és egy végleges 1000-utas softmax (kimeneti) réteg követ. A technológiáról és annak képelemző képességeiről bővebben: Krizhevsky, A., Sutskever, I. és Geoffrey Hinton (2012): ImageNet classification with deep convolutional neural networks. *Communications of the ACM*. 60, 84-90.o. <https://doi.org/10.1145/3065386>

hálózatok népszerűsége évről évre nőtt, ma már számos területet érintenek és látványos fejlődést hoztak az MI fejlesztésekben, különösen olyan alkalmazásokban, mint a gépi látás, orvosdiagnosztika, beszédtechnológia és hangfeldolgozás képfelismerés, természetes nyelvfeldolgozás és az önvezető rendszerek.

A mélytanulás mellett, a leginkább elterjedt gépi tanulási mintázatok közé a felügyelt, felügyelet nélküli és megerősítéses gépi tanulás tartoznak.

A felügyelt tanulásra (supervised learning) jellemző, hogy a modellt előre címkézett adatokkal (meghatározott input-output párokkal) képzik. Mivel a tanulóadatokhoz előzetesen hozzá van rendelve egy adott érték, a modell idővel megtanulja, hogyan hozza létre az ahhoz kötődő helyes kimenetet. Elérendő célja az általánosítási képesség elsajátítása, tehát, hogy elegendő minta feldolgozását követően a modell képes legyen olyan adat-címke párok esetében is helyes döntést hozni, amelyeket korábban nem ismert.³¹ Példaként szolgálhat egy e-mail szűrő algoritmus, amely spam és nem spam üzenetek között tud különbséget tenni. Ebben az esetben az e-maileket előre megjelölik "spam" vagy "nem spam" címkével. Ahogy az algoritmus több és több e-mailt dolgoz fel, megtanulja felismerni a spam e-mailek jellemzőit, még olyan új üzenetek esetében is, amelyekkel korábban nem találkozott. így válva fokozatosan hatékonyabbá a nem kívánt e-mailek szűrésében.

A felügyelet nélküli tanulásnál (unsupervised learning) a modell előzetesen nem címkézett adatokból tanul, melyekhez nincs hozzárendelve előre meghatározott érték vagy megfelelő kimenet-pár. Itt a modellnek önállóan kell felfedeznie a rejtett struktúrákat, mintázatok a bemeneti adatokban.³² A tanulás során felfedezett összefüggések később sokféleképpen hasznosíthatók, az ilyen modellek egyik fő alkalmazási területe az ún. klaszterezés, ahol az algoritmus az adatpontokat hasonlóságuk alapján meghatározott csoportokba rendezi. Például egy ügyfélszegmentációs eljárásnál a modell az ügyfelek vásárlási szokásai, demográfiai adatai és egyéb jellemzői alapján csoportokat hoz létre mely által azonosíthatóvá válnak az ügyfelek különböző szegmensei (gyakori vásárlók, kisebb megtakarítással rendelkező ügyfelek, prémium vásárlók stb.).³³

³¹ Hastie, T., Tibshirani, R., Friedman, J. (2009): Overview of Supervised Learning. In: The Elements of Statistical Learning. Springer Series in Statistics. Springer, New York. 9-41.o.

³² Hinton, Geoffrey – Terrence Sejnowski J. (Szerk.) (1999): Unsupervised Learning: Foundations of Neural Computation. Bradford Books, The MIT Press. 4-19.o

³³ E két eljárás ötvözetét alkotja egy újabb megközelítés, az ún. önfelügyelt gépi tanulás (self-supervised learning), mely során az algoritmusok saját magukat képezik ki a rendelkezésre álló adatokon keresztül, anélkül, hogy szükség lenne ember által előre címkézett adatokra vagy felügyeletre. Azért nevezhető az előzőleg említett két

A felügyelet nélküli tanuláshoz hasonlóan, a megerősítéses tanulásnál (reinforcement learning) sem meghatározott kimeneti értékkel bír, előre címkézett adatokat használ a modell. Ehelyett feladata, hogy saját maga következtessen ki, hogy miként oldható meg optimálisan egy adott feladatot. Ezt úgy éri el, hogy különböző műveleteket hajt végre - eközben az algoritmus (agent) próbálkozik és hibázik - majd a képzés során a konkrét döntésekhez pozitív vagy negatív visszajelzést (jutalmakat vagy büntetéseket) csatolnak. A modell célja, hogy maximalizálja a jutalmakat és minimalizálja a büntetéseket.³⁴ A megerősítéses tanulás elsősorban a dinamikus környezetekben való döntéshozatalra összpontosít, gyakori alkalmazási területei robotika, a táblajátékok (AlphaGo is ilyen módon lett fejlesztve) vagy az önvezető autók. Az utóbbi esetben a modell megtanulja, hogy milyen manővereket hajtson végre bizonyos forgalmi helyzetekben, például mikor kell megállnia egy piros lámpánál vagy elkerülnie egy akadályt. Mindezt úgy teszi, hogy kipróbál különböző stratégiákat, és a következmények alapján, például egy ütközés vagy sikeres manőver, jutalmat vagy büntetést kap.

Az utóbbi években vált egyre inkább kiemelt témává a gépi tanulás kutatásaiban az un. informált gépi tanulás (informed machine learning).³⁵ Ez a megközelítés lehetővé teszi, hogy a modellek előzetesen meglévő tudást, fizikai összefüggéseket, szabályokat és korlátozásokat építsenek be a tanulási folyamatba. Míg a korábban tárgyalt eljárások - a hagyományos, tisztán adatvezérelt modellek - próbálgatás útján tanulnak az adatokból, az informált gépi tanulás során a fejlesztők bizonyos előzetes szabályokat is megadnak a modellnek, hogy segítsék a tanulást, így azok mind szabályokra, mind adatokra támaszkodnak a tanulási folyamatban. Ez a technika több tekintetben is hasznos lehet: (1) Előfordulhat, hogy nem áll rendelkezésre elegendő adat a

tanulási eljárás keverékének, mert egyrészt a felügyelet nélküli tanításhoz hasonlóan lehetővé teszi, hogy a modellek nagy mennyiségű címkézetlen adatból tanuljanak, másrészt kihasználja a felügyelt tanulás hatékonyságát az algoritmusok által mesterségesen generált címkék alkalmazásával. Amellett, hogy csökkenti a címkézési erőfeszítéseket és költségeket, egyúttal képes mélyebb reprezentációk és általánosító képességek fejlesztésére. Ilyen gépi tanulást alkalmaz például a Meta AI által fejlesztett data2vec algoritmus is, amit a gépi látás és a későbbiekben még részletesen tárgyalt Nagy Nyelvi Modellek fejlesztésére használnak. Lásd bővebben: Kejriwa, Kunal (2023): Data2vec: A Milestone in Self-Supervised Learning. UniteAI. Elérhető: <https://www.unite.ai/data2vec> (2023. 11. 04.)

³⁴ Mnih, Volodymyr, Kavukcuoglu, K., Silver, D., Rusu, A., Veness, J., Bellemare, M., Graves, A., Riedmiller, M., Fidjeland, A., Ostrovski, G., Petersen, S., Beattie, C., Sadik, A., Antonoglou, I., King, H., Kumaran, D., Wierstra, D., Legg, S., & Hassabis, D., (2015): Human-level control through deep reinforcement learning. Nature, 529-533.o.

³⁵ Bár az informált gépi tanulás ötlete nem új, de az utóbbi években vált egyre inkább népszerűvé és kezdett elterjedni a gépi tanulás kutatásaiban. Az alapjául szolgáló ötlet (modellek kiegészítése emberi ismeretekkel) már korábban is jelen volt bizonyos formákban, például a szakértői rendszerekben az 1980-as és 1990-es években. Az adatvezérelt megközelítések és a nagy adatkészletek elérhetősége miatti előretörés után a kutatók rájöttek, hogy a tisztán adatvezérelt modellek korlátai, például az adathalmazok torzítottsága vagy a hiányos adatok, megkövetelik az előzetes tudás integrálását a modellekbe. Ez különösen akkor vált fontossá, amikor a modelleknek olyan területeken kell helyesen működniük, ahol a pontosság, a megbízhatóság és a szabályozási megfelelés kritikus, mint például az egészségügy, a pénzügyek és az autonóm járművek.

jól teljesítő és megfelelően általánosított modellek betanítására. Az informált gépi tanulás segíthet kiegészíteni a hiányzó adatokat; (2) A pusztán adatokon alapuló modellek figyelmen kívül hagyhatják a természeti törvényeket, szabályozásokat vagy biztonsági irányelveket, amelyek a megbízható MI megvalósításához elengedhetetlenek. Az informált tanulás beépítheti ezeket a korlátozásokat (3) A tisztán adatvezérelt modellek visszatükrözhetik az adatok torzításait és elfogultságait. Az előzetes információk és korlátozások segíthetnek kiküszöbölni ezeket a problémákat. (4) Az informált tanulással könnyebben átláthatóvá tehetők a modellek, ami a megbízhatóság növelése mellett a szabályozói követelmények teljesítését is elősegíti. Nem utolsó sorban pedig a kiegészítő információk stabilabbá és robusztusabbá tehetik a modelleket, ami különösen a korlátozott, zavaros vagy egymásnak ellentmondó tanítóadatok esetén fontos.³⁶

A következő fejezetben e technológiák konkrét alkalmazási területeinek bemutatására és az ezekhez kapcsolódó lehetséges veszélyek feltérképezésére kerül sor.

³⁶ von Rueden, L., Mayer, S., Beckh, K., Georgiev, B., Giesselbach, S., Heese, R., Kirsch, B., Pfrommer, J., Pick, A., Ramamurthy, R., Walczak, M., Garcke, J., Bauckhage, C. and Schuecker, J., (2019): Informed Machine Learning – A Taxonomy and Survey of Integrating Prior Knowledge into Learning Systems. IEEE Transactions on Knowledge and Data Engineering, 35, 614-633.o.

3 Alkalmazások és kihívások

3.1 MI a köz szolgálatában?

Az MI térnyerése az élet számos területen érezteti hatását. Az elmúlt évtizedben egyre gyakrabban alkalmaztak gépi eljárásokat a mindennapi életet közvetlenül befolyásoló, kritikus döntések meghozatalakor: MI-alapú diagnosztikai eszközök segítenek az orvosoknak a daganatos betegségek gyorsabb felismerésében; MI-számításokra támaszkodnak a bankok a hitelkérelmek automatikus kiértékelésekor; a biztosítóintézetek az ügyfélkockázati minősítések elkészítésekor; az adóhatóságok, az állampolgárok adóelkerülési gyakorlatának azonosítására; a szociális ellátórendszerek, hogy megállapítsák egy adott személy állami ellátásokra és szolgáltatásokra való jogosultságát, de még a bűnüldöző hatóságok is annak érdekében, hogy pontosabb előrejelzéseket készíthessenek a jövőbeni bűnelkövetés vagy bűnismétlés valószínűségéről. E rendszerek alkalmazási területei között megkülönböztetett figyelmet érdemelnek a hatóságok által közigazgatásban használt MI technológiák, hiszen az ebben a környezetben hozott döntések érdemi hatással lehetnek az állampolgárok jogaira és kötelezettségeire is például a közellátások és közszolgáltatások, a lakhatás, foglalkoztatás, oktatás vagy az igazságszolgáltatás területén. A következőkben néhány már napjainkban is alkalmazott gyakorlat bemutatására kerül sor.

3.1.1 RPA és ADM - az automatizált döntéshozó rendszerek

Ahogy a döntéshozatali folyamatokban egyre nagyobb hangsúly összpontosul az adatokra, az azokat feldolgozni, kategorizálni és értékelni képes eljárások is egyre növekvő részét teszik ki az adminisztratív gyakorlatoknak.³⁷ A rendelkezésre álló erőforrások elosztásának optimalizálása, a közszolgáltatások személyre szabása, a közösségi részvétel elősegítése vagy

³⁷ A kormányok világszerte igyekeznek kihasználni az MI és Big Data technológiák előnyeit azáltal, hogy saját adatbázisaikat összekötve, rendszerezett formában gyűjtik, tárolják és dolgozzák fel a rendelkezésükre álló adatokat. Az általuk birtokolt adatok segítségével pontosabb képet kaphatnak az állampolgárok életéről, a kormányzati szervek működéséről és a központi, illetve regionális közigazgatási viszonyokról is. Általánosságban kijelenthető, hogy e technológiák elterjedése és az adatbázisok összekötése lehetőséget teremt a közigazgatás széles körű modernizálására, olcsóbban, eredményesebben és átláthatóbban működő közszolgáltatások biztosítására, ezeken keresztül pedig az állampolgárok életkörülményeinek általános javítására is. Lásd bővebben: Maciejewski, Mariusz (2016): To do more, better, faster and more cheaply: using big data in public administration. *International Review of Administrative Sciences*, 83, 120–135.o.

az okosvárosokkal kapcsolatos adatvezérelt tervezés csak néhány azon a felhasználási területekből melyekben ilyen rendszereket alkalmaznak ma is.³⁸

A Robotizált Folyamat-Automatizálás (Robotic Process Automation, továbbiakban RPA) révén lehetségessé válik az ismétlődő, előre meghatározott és kiszámítható lépésekből álló adminisztratív feladatok automatizálása. Jelenleg az RPA rendszereket leginkább olyan kevésbé összetett adminisztratív feladatok automatizálása használják, mint az adatrögzítés-és változtatás, űrlapok kitöltése és frissítése, vagy automatikus emlékeztetők és értesítők kiküldése.³⁹

Abban az esetben azonban, amikor az RPA technológiát más MI rendszerekkel - például természetes nyelvfeldolgozási modellekkel – ötvözik, lényegesen fejlettebb Automatizált Döntéshozó Rendszerek (Automated Decision-Making Systems, ADM) létrehozása válik lehetségessé. Az ilyen rendszereket többek között a közigazgatási hivatalokba nagyszámban beérkező engedélykérelmek kezelésére, valamint különböző pénzügyi/számviteli feladatok (pl. számlák elkészítése, a kifizetések nyomon követése pénzügyi jelentések készítése vagy bérszámfejtés) elvégzésére lehet használni. Előnyei közé tartozik, hogy képes autonóm módon frissíteni nagy és dinamikusan változó adathalmazokat (pl. az állampolgárok személyi adatainak és lakcímének nyilvántartását) valamint, hogy különböző anomáliákat tud azonosítani strukturált információs közegekben. Ennek köszönhetően az ilyen technológiák gyakran kerülnek alkalmazásra a szociális ellátások rendszerében, ahol a különböző állami segítségnyújtási ellátásokra és szolgáltatásokra való jogosultságának értékelésére, valamint az ilyen ellátások és szolgáltatások nyújtásával, csökkentésével visszavonásával vagy visszaigénylésével kapcsolatos döntések meghozatalában játszanak szerepet.⁴⁰

A svédországi Trelleborg városa például már évek óta használ RPA rendszert annak érdekében, hogy gyorsabbá és költséghatékonyabbá tegye a szociális segélyezés ügymenetét. A nagyszámban beérkező kérelmek elbírálásával kapcsolatos nehézségek miatt – ilyen volt a

³⁸ Misuraca, Gianluca., és van Noordt, C. (2020): Overview of the use and impact of AI in public services in the EU. EUR 30255 EN, Publications Office of the European Union, Luxembourg, 39-52.o.

³⁹ Houy, C., Hamberg, M. & Fettke, P., (2019): Robotic Process Automation in Public Administrations. In: Räckers, M., Halsbenning, S., Rätz, D., Richter, D. & Schweighofer, E. (Hrsg.), Digitalisierung von Staat und Verwaltung. Bonn: Gesellschaft für Informatik. 62-74.

⁴⁰ Ranerup, Agneta & Henriksen, Z. H (2019): Value positions viewed through the lens of automated decision-making: The case of social services. Government Information Quarterly. Volume 36(4) 2-3.o.

késedelmes kifizetés - még 2016-ban döntött úgy Trelleborg önkormányzata, hogy digitalizálja és automatizálja a szociális ellátások adminisztrációs rendszerét. A trelleborgi modellt elsősorban arra tervezték, hogy a különböző szociális segélykérelmek (pl. otthoni ápolás, a betegségi ellátás, a munkanélküli segély) feldolgozásában és elbírálásában segítsen a hivataloknak. A modernizáció sikerét mutatja, hogy bevezetést követően a szociális támogatásra irányuló digitális kérelmek 85%-át ma már RPA rendszerek kezelik. ⁴¹

3.2 Prediktív analitika az adatvezérelt döntéshozatal szolgálatában

A prediktív technológiát alkalmazó MI rendszerek kiterjedt adathalmazok elemzésével különböző mintákat és összefüggéseket képesek feltárni, melyekre alapozva később megbízható előrejelzéseket készítenek a jövőt illetően. A döntéstámogatásban betöltött szerepük jelentős, hiszen segítségükkel jobb minőségű szakpolitikai kormányprogramok-és stratégiák megalkotása válik lehetségessé. A prediktív analitika technológiák döntéstámogató alkalmazására jó példaként szolgálhat a 2019-ben TalTecCity néven indított észt kezdeményezés, mely intelligens szenzorok segítségével gyűjt információkat egyebek mellett az észt főváros (Tallinn) területén belül közlekedő emberek és járművek mozgásáról, az utak és parkolók forgalmáról, a levegőben mérhető külső légszennyezésről, és a köztéren elhelyezett hulladékártalók telítettségéről. Az így nyert információkra alapozva később modellek hoznak létre, amelyek nagyban segítenek az önkormányzati döntéshozóknak abban, hogy adatvezérelt szakpolitikai stratégiákat dolgozzanak ki a várostervezéssel, a közlekedéssel és a környezetgazdálkodással kapcsolatos kérdésekben. ⁴²

3.2.1 Adózással és szociális juttatásokkal kapcsolatos csalások

Az elmúlt években egyre gyakrabban fordul elő, hogy a nemzeti adóhatóságok is MI alapú előrejelző rendszereket alkalmaznak a pénzügyi visszaélések hatékonyabb feltérképezése céljából. Az ilyen rendszerek nem csupán a folyamatban lévő illegális tevékenységeket képesek észlelni, hanem adat- és kockázatelemző algoritmusok segítségével a jövőbeni csalások valószínűségét is előre jelzik.

⁴¹ Ranerup, Agneta, Henriksen, Z. H. (2022): Digital Discretion: Unpacking Human and Technological Agency in Automated Decision Making in Sweden's Social Services. *Social Science Computer Review*. 40(2), 445-461.o.

⁴² FRA Report (2020): European Union Agency for Fundamental Rights (2020): Artificial Intelligence, Big Data and Fundamental Rights Country - Research Estonia 2020.

A magyar Nemzeti Adó- és Vámhivatal is hasonló MI alapú előrejelzéseket használ az adócsalásra és adóelkerülésre irányuló gyakorlatok felderítésére. A Rugalmas Adóellenőrzési Döntéstámogató és Adatbányászati Rendszert (RADAR) a korábban vizsgált ügyek eredményei alapján olyan jellemzőket és ismérveket keres, amelyek nagy valószínűséggel vezettek magas adóhiányhoz a múltban. Az eredmények kiértékelésével megalapozott előrejelzéseket és következtetéseket tud készíteni a lehetséges jövőbeli esetekre vonatkozóan.⁴³

A prediktív analitika alkalmazására további példaként szolgálhat a Hollandiában 2020-ig használt SyrRy (Systeem Risico Indicatie) kockázatértékelő rendszer, melyet a potenciális jövőbeli társadalombiztosítási, szociális vagy egyéb jóléti juttatással kapcsolatos csalások felderítésére használtak. Az eszköz a meglévő kormányzati adatállományokból (pl. adóbevallások, társadalombiztosítási adatok, egészségügyi kórtörténet) származó kockázati mutatókat elemezte annak érdekében, hogy azonosítsa azokat a személyeket, akiknél fokozottan fennáll a jóléti juttatásokkal való csalás vagy visszaélés kockázata. Érdeemes megjegyezni azonban e rendszerrel kapcsolatban, hogy bevezetése után nem sokkal több jogvédő szervezet is kritikával illette annak működését. A kritikák tárgyát képezte, hogy a SyrRy rendszer több alkalommal is megsértette a magánélethez fűződő alapvető jogokat, és hátrányosan diszkriminálta a gazdaságilag sebezhető állampolgárokat. További aggodalomként vetették fel a rendszer belső működésének átláthatatlan voltát, valamint azt, hogy a döntésben érintett személyeknek nem volt lehetőségük megismerni a döntéseket megalapozó adatokat. A szoftver alkalmazását - a holland bíróság döntése nyomán – végül 2020-ban hivatalosan is leállították.⁴⁴

3.2.2 Rendvédelmi szervek és büntető igazságszolgáltatás

A prediktív technológiák talán leginkább megosztó felhasználási területe az igazságszolgáltatási rendszerben történő alkalmazásukhoz kötődik. Az e területeken jellemzően profilozáson, pontozáson, kockázatelemzésen és magatartáselőrejelzésen alapuló algoritmusokra rengeteg figyelem összpontosult az elmúlt években, nem egy esetben heves kritikák keresttüzébe állítván őket. Szerepük és befolyásuk nem elhanyagolható hiszen az általuk készített értékelések befolyásolhatják a valós életben alkalmazott bűnmegelőzési

⁴³Fejes Erzsébet - Futó Iván (2021): Mesterséges intelligencia a közigazgatásban – az érdemi ügyintézés támogatása. Pénzügyi Szemle. Különszám 2021/1, 44.o.

⁴⁴ Misuraca, Gianluca., - van Noordt, C. (2020): Overview of the use and impact of AI in public services in the EU. EUR 30255 EN, Publications Office of the European Union, Luxembourg 45-46.o.

stratégiákat, többek között azt, hogy hol szükséges fokozni a megfigyelést vagy mely személyek igazoltatása válhat indokolttá. Emellett a büntetőügyekben hozott bírósági döntésekre is hatással lehetnek, például az őrizetbe vétellel, az óvadék megállapításával, az előzetes letartóztatással, a próbára bocsátással vagy a büntetés-végrehajtás felfüggesztésével kapcsolatosan. ⁴⁵

A bűnüldöző hatóságok esetében az előrejelző technológiák egyik legelterjedtebb felhasználási módja az úgynevezett prediktív rendfenntartáshoz (Predictive Policing) kötődik. E gyakorlat a rendelkezésre álló bűnügyi adatok feldolgozásával statisztikai módszereket alkalmaz annak előrejelzésére, hogy hol és mely időszakokban magasabb a jövőbeni bűncselekmények elkövetésének valószínűsége. Ilyen előrejelző eszközöket használnak a bűnüldöző szervek például Hollandiában (ProKid), Svájcban (Precops), az Egyesült Királyságban (NDAS), Németországban (RADAR-iTE) és Olaszországban (Delia) is. ⁴⁶

A bűnüldöző hatóságok által használt előrejelzések másik gyakori felhasználási módjára példaként szolgálhat az Egyesült Királyságban évek óta használt HART kockázatelemző-rendszer (Harm Assessment Risk Tool). A gépi tanuláson alapuló szoftvert a durhami rendőrség használja a jövőbeni bűnismétlés valószínűségének előrejelzésére. A HART rendszer – 34 kategóriát magába foglaló adatbázisból - profilt alkot az eljárás alá vont személyről, majd annak kiértékelésével egy kockázati pontszámokat hoz létre. A kapott pontszám alapján a vizsgált személyt a három előre meghatározott kockázati csoport (magas, közepes, alacsony) valamelyikébe sorolja. Ez az osztályozás segíti az igazságügyi hatóságok döntését a tekintetben, hogy vád alá helyezték-e a személyt, vagy a Checkpoint névre hallgató, bíróságon kívüli rehabilitációs programba utalják. Amennyiben az elkövető az alacsony kockázati csoportba kerül, a programba utalják, ezzel lehetőséget biztosítva számára, hogy elkerülje a bírósági eljárás lefolytatását. Ellenkező esetben – tehát ha közepes, vagy magas kockázati csoportba sorolták - megindul a vádemelés. ⁴⁷

⁴⁵ Herke Csongor (2021): A mesterséges intelligencia kriminalisztikai aspektusai. *Belügyi Szemle*. 69/10. 1714–1720.o.

⁴⁶ Fair Trials (2021): Automating Injustice - The Use of Artificial Intelligence & Automated Decision- Making Systems in Criminal Justice. 8-21.o.; Fantoly Zsanett, Herke Csongor, Szabó Barbara (2023): The role of AI-based systems in negotiated proceedings. *e-Revue Internationale de Droit Pénal*. 2023, A-18, 4.o.

⁴⁷ Oswald et al (2018): Algorithmic risk assessment policing models: lessons from the Durham HART model and 'Experimental' proportionality. *Information & Communications Technology Law*. 27(2), 223-250.o.

Az előzőekben ismertetett példák csak egy részét képezik a közigazgatási MI gyakorlatok egyre növekvő körének. Mivel az ehhez hatással vannak számos, az emberek életét közvetlenül is befolyásoló kérdésre - például, hogy milyen típusú és minőségű közszolgáltatásokhoz férhetnek hozzá, vagy milyen bánásmódban részesülnek igazságszolgáltatási eljárások során - szükséges megérteni és átlátni azt, hogy melyek a gyengeségei és veszélyei az ilyen technológiáknak, a következő fejezet erre tesz kísérletet.

3.3 Az MI alkalmazásból eredő kihívások, kockázatok

Az elmúlt években az MI-rendszerek alkalmazásából eredő társadalmi, etikai és jogi kockázatokkal kapcsolatban a figyelem elsősorban három kulcsfontosságú kihívásra irányult: (1) az MI-re jellemző gépi tanulásból és autonóm működésből eredő átláthatóság, elszámoltathatóság és kiszámíthatóság hiánya; (2) az MI által vezérelt téves, elfogult vagy diszkriminatív döntéshozatal veszélye; (3) az MI-rendszerek alapvető emberi szabadságjogokra és a demokratikus intézmények működésére gyakorolt negatív hatása.⁴⁸

3.3.1 A Black Box probléma – átláthatóság, értelmezhetőség, indoklási kötelezettség

Az MI alapú döntéshozatalra jellemző átláthatatlanság – Jenna Burrell felosztása alapján – három alkategóriára bontható. Szándékos (intentional) az átláthatatlanság, ha az algoritmikus döntéshozatal mögött álló folyamatokat tudatosan tartják rejtve (ilyen lehet például a Google keresőmotorjának programkódja, mely üzleti titkoknak minősül). Az írástudatlan (illiterate) átláthatatlanság azt jelenti, hogy ugyan senki nem akadályozza egy adott algoritmus belső működésének megismerését, mégis annak rendkívül összetettsége okán azt csak nagyon kevesen – elsősorban a programnyelvet ismerő, szakismerettel bíró személyek – tudják értelmezni. A lényegbeli (intrinsic) átláthatatlanság pedig azt jelenti, hogy egyes rendszerek olyan dinamikus – sok esetben autonóm módon – változnak, hogy még a magasan képzett programozóknak – vagy a programkód megalkotóinak – is kihívást jelent teljes egészében megérteni alapvető működésüket.⁴⁹

⁴⁸ Az MI társadalmi, etikai és jogi kihívásaival kapcsolatban gazdag magyar szakirodalom áll rendelkezésre. Itt csak az utóbbi években megjelent könyvekre utalva: Zódi Zsolt (2018): *Platformok, robotok és a jog. Új szabályozási kihívások az információs társadalomban*. Budapest, Gondolat Kiadó; Héder Mihály (2020): *Mesterséges Intelligencia. Filozófiai kérdések, gyakorlati válaszok*. Budapest, Gondolat Kiadó; Tilesch György – Hatamleh, Omar (2021): *Mesterség és Intelligencia. Vegyük a kezünkbe a sorsunkat az MI korában*. Budapest, Libri Kiadó; Csepeli György (2020): *Ember 2.0 – A mesterséges intelligencia gazdasági és társadalmi hatásai*. Budapest, Kossuth Kiadó; Török Bernát - Zódi Zsolt (szerk.) (2021): *A mesterséges intelligencia szabályozási kihívásai. Tanulmányok a mesterséges intelligencia és a jog határterületeiről*. Budapest, Ludovika Egyetemi Kiadó.

⁴⁹ Burrell, Jenna (2016): *How the Machine 'Thinks: Understanding Opacity in Machine Learning Algorithms*. *Big Data and Society*. 3/1. 3-6.o.

Lényegében ezt az utolsó esetet írja körbe a gépi döntéshozatallal kapcsolatosan gyakran említett Black Box probléma is (fekete-doboz probléma), mely a fejlett, gépi tanuláson alapuló MI modellekhez köthető legsürgetőbb kihívások egyike. Ezek az összetett algoritmusok által vezérelt, önállóan tanulni képes MI-k olyan módon juthatnak egy adott következtetésre, amelyet az azt alkalmazók nem minden esetben láthatnak előre.⁵⁰ A végeredményt ugyan képesek értelmezni, annak módját azonban, hogy milyen összefüggéseken keresztül jutott el a program a válaszig, már nem látják át teljesen (máshogy kifejezve: értik a mit, de nem értik meg a miértet). Bár nem használja a Black Box kifejezést, erre a jelenségre utal 2017-ben megjelent Pszichopolitika című könyvében a német-koreai filozófus Byung-Chul Han is, amikor az *adattotalitarizmus* és adatalapú gépi döntések elterjedésének következményeiről ír, mely *a tudás új korszakát jelenti be. A korrelációk helyettesítik az okozatiságot. Az így van-ez helyettesíti a hogyant.*⁵¹

Az átláthatatlan működésből következik az MI-alapú gépi döntések értelmezhetőségének és magyarázhatóságának (interpretability/ explainability) kettős problémája is.⁵² Bár e két fogalmat sok esetben szinonimaként használják, hiszen hasonló célt szolgálnak - az MI alapú döntések és előrejelzések megindoklását – alapvető különbség van a jelentéstartalmuk között. Az értelmezhetőség a bemenetek és kimenetek közötti ok-okozati összefüggések emberi megértésének mértékét jelenti, ami magában foglalja, annak átlátását, hogy a modell milyen döntési szabályokat alkalmaz és hogy mely jellemzők milyen fontossággal bírnak a kimenetek elérésében. Egyúttal jelenti azt is, hogy a rendszer fejlesztői vagy alkalmazói képesek-e előre

⁵⁰ Pasquale, Frank. (2015): The Black Box Society: The Secret Algorithms That Control Money and Information. Harvard University Press.

⁵¹ Byung megfogalmazásában: *Azzal párhuzamosan azonban, hogy a társadalmak korábban elképzelhetetlen mennyiségű adatot halmaznak fel az őket körülvevő világgal kapcsolatban, az információfeldolgozás és -megértés hagyományos módzatai is lényeges változáson mennek keresztül. Korunkban, amit az adattotalitarizmus és adatfétisizmus lelkesít át, a Big Data alapú technológiák jelentik az iránytűt a biztos tudás felé. Minden mérhető és számszerűsíthető. A dolgok elveszítik mostanáig rejtve maradt, titkos összefüggéseiket.* Lásd: Han, Byung-Chul (2020): Pszichopolitika. (Ford. Csordás Gábor) Budapest, Typotex. 91.o.

⁵² Ezzel kapcsolatban kiemelendő, hogy az értelmezhetőség - magyarázhatóság kettős problémája arra is rámutat, hogy nem egyértelmű, hogy mit, kinek és hogyan kell magyarázni. Először is, mit kell magyarázni? Az MI-rendszerek működési mechanizmusait, az általuk hozott döntéseket és előrejelzéseket? Esetleg a tervezési folyamatot és az alkalmazási célokat? Mindez egyaránt fontos lehet, de eltérő magyarázati megközelítést igényel. Másodszor, kinek kell magyarázni? Az alkalmazónak, a felhasználóknak, fejlesztőknek, az érintett személyeknek, a szabályozó hatóságoknak? Minden érintett csoport más-más szintű és típusú magyarázatra lehet kíváncsi, más-más háttértudással rendelkezik. Harmadszor, ki adja meg a magyarázatokat? Maga az MI-rendszer? A fejlesztők? Egy független szakértő? Attól függően, hogy ki magyaráz, a megközelítés is eltérő lehet. Világos, hogy minden magyarázat egyfajta "fordítást" jelent a célközönség számára. Vagyis a magyarázatokat mindig a konkrét felhasználói helyzethez és igényekhez kell igazítani, hogy érthetőek és hasznosak legyenek (egy szakértőnek szóló technikai magyarázat például más lesz, mint egy laikus számára készült leírás. Lásd a (meg)magyarázható MI-ről (explainable AI, XAI) bővebben O'Hara, K. (2020): Explainable AI and the philosophy and practice of explanation. Computer Law & Security Review, Vol.39. 3-6.o.

jelezni, mi történik, ha változtatásokat hajtanak végre a bemeneti adatokon vagy az algoritmus paraméterein. A magyarázhatóság ezzel szemben az MI-rendszerek döntéseinek és előrejelzéseinek érthető magyarázataira összpontosít, ideértve a modell kimenetének okait és indoklását. Míg az értelmezhetőség a modell belső mechanizmusaira vonatkozik, addig a magyarázhatóság túlmutat a modell belső működésén, és arra irányul, hogy a rendszerek által hozott döntéseket és előrejelzéseket érthető módon közvetítse a felhasználóknak és érintetteknek.⁵³ Tehát az értelmezhetőség a modell "mikéntjére" összpontosít, a magyarázhatóság pedig a "miértjére" ad választ.

Mind az értelmezhetőség mind pedig a magyarázhatóság hiánya különösen aggasztó lehet azon helyzetekben, amikor az algoritmusok olyan kritikus döntéseket hoznak melyek lényeges társadalmi vagy jogi következményekkel járhatnak az egyénre nézve, mint például az előzőleg ismertetett közigazgatási hatóságok, bűnüldöző és rendvédelmi szervek által alkalmazott rendszerek esetében. Ilyenkor a tisztességes eljáráshoz való jog részeként merül fel a hatósági szerv oldalán jelentkező indoklási kötelezettség, melynek elmulasztása csökkentheti a döntések elfogadhatóságát, és az állampolgárok az önkéntes jogkövetési hajlandóságát.⁵⁴ Jogosan merül fel tehát az az igény, hogy MI-rendszerek széles körben elfogadottá és megbízhatóvá válásához, elengedhetetlen az átláthatóság - és az ahhoz szorosan köthető elszámoltathatóság biztosítása. Ez magában foglalja az algoritmusok és döntéshozatali folyamatok érthetővé és magyarázhatóvá tételét (indoklási kötelezettség) valamint annak lehetőségét, hogy az érintettek megkérdőjelezhessék vagy felülvizsgálhassák a rendszer által hozott döntéseket (ha szükséges emberi felügyelet kikövetelésével).⁵⁵

Amennyiben ez nem történne meg és az adattotalitarizmusból eredő új értelmezési keretek kerülnének előtérbe, ahol a "miért" kérdése helyett a "mit" válik iránytűvé, és ahol az összefüggések helyett csupán száraz adatokra támaszkodunk, új előre nehezen megjósolható

⁵³ Rudin, C. (2019): Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nat Mach Intell.* (1), 206–215.o; Molnar, Christoph (2020): *Interpretable Machine Learning: A Guide for Making Black Box Models Explainable.* Leanpub. Molnar a könyvet eredeti nyelven githubon szabadon elérhetővé tette: <https://christophm.github.io/interpretable-ml-book/>

⁵⁴ Hohmann Balázs (2021): A mesterséges intelligencia közigazgatási hatósági eljárásban való alkalmazhatósága a tisztességes eljáráshoz való jog tükrében. In Török Bernát és Zódi Zsolt (szerk) *A mesterséges intelligencia szabályozási kihívásai.* Budapest, Ludovika Egyetemi kiadó. 413.o.

⁵⁵ Elméletben az emberi felügyeletet lehetőségét garantálja a GDPR 22. cikk (1) bekezdése, mely kimondja, hogy az adatalany jogosult arra, hogy ne terjedjen ki rá az olyan, kizárólag automatizált adatkezelésen alapuló döntés hatálya, amely rá nézve joghatással járna vagy őt jelentős mértékben érintené. A gyakorlatban azonban a gépi döntések végrehajtásakor az algoritmus alkalmazása mellett jellemzően megmarad az emberi tényező is, vagyis a döntés nem kizárólag automatizált módon történik, így gyakran nem terjed ki a jogszabály védelme a döntés alanyára. (Például a banki hitelbírálat esetén az algoritmus dönt az ügyfél hitelképességéről, de a döntés végső hitelesítését egy banki alkalmazott végzi el).

kihívásokkal találhatja szemben magát a jövő társadalma. A Big Data technológiákon alapuló előre jelző MI-k kapcsán Zödi Zsolt Platformok, robotok és a jog című könyvében szellemes illusztrálja az átláthatatlanságból eredő potenciális jogi és morális dilemmákat:

„Képzeljük el azt a helyzetet, hogy mondjuk bebizonyosodik, hogy a 47-es lábú fekete hajú, 195 centiméter magas, 19 éves fiatal férfiak nagyon nagy eséllyel követnek el erőszakos bűncselekményeket. (Szándékosan írtam ilyen abszurd paramétereket.) Az persze elképzelhetetlen, hogy preventív céllal ezeknek az embereknek a személyes szabadságát korlátozzuk, de egy idő után nagy lesz a csábítás, hogy legalább valamilyen módon például megfigyeljük őket. De még a megfigyelés sem lesz a hagyományos igazolási technikákkal igazolható. Hogyan magyarázható, hogy az intézkedés tartalmát semmilyen szempontból nem érintő paraméter az intézkedés indoka? Ez a gondolkodásmód teljes szakítás lenne a felvilágosodás eszményeivel; például azzal a megfontolással, hogy csak olyanért vagyok felelőssé tehető, amelyen magam is képes vagyok változtatni.”⁵⁶

3.3.2 Egyenlő bánásmód követelményének megsértése

Az MI-alapú gépi döntéshozatal kapcsán gyakran elhangzó érv, hogy az ilyen eljárások képesek kizárni az emberi elfogultságot és előítéleteket a döntéshozatali folyamatok során. A valóságban azonban ez csak annyira érhető el, amennyire az e rendszereket működtető programkódok, illetve az e rendszereket tanító adatkészletek is pártatlanok. Ha az MI elfogult, torzított adathalmazokon tanul, akkor hiányos, helytelen vagy elfogult képet kap a valóságról

Az elemző és előre jelző technológiák egyik legfőbb haszna abban rejlik, hogy a hatalmas adathalmazok begyűjtésével, feldolgozásával és kiértékelésével korábban nem realizált minták és összefüggések feltárása válik általuk lehetségessé. Azonban pontosan ennek okán fordulhat elő az, hogy bizonyos minták és összefüggések algoritmusok által történő felismerése, majd ezek szabályszerűségnek értékelése valamely, a társadalomban már korábban jelenlévő előítélet vagy elfogultság akaratlan legitimálásához vezethet.

Az automatikus döntéshozatal kapcsán az elfogultság három fő típusát lehet elkülöníteni; (1) szándékos, (2) nem szándékos és a (3) tanult elfogultság. Szándékos az elfogultság, ha a rendszert fejlesztője, vagy annak alkalmazója tudatosan épít be előítéleteket a döntéshozó

⁵⁶ Zödi Zsolt (2018): Platformok, robotok és a jog. Új szabályozási kihívások az információs társadalomban. Budapest, Gondolat Kiadó. 234-235.o.

algoritmusba (például gazdasági érdekekből). Ide tartoznak az árdiszkriminációs gyakorlatok és a dinamikus árazás (differenciált árképzés) bizonyos esetei. E marketing stratégiák során az online platformok potenciális vásárlói eltérő árat látnak ugyanazon termékért vagy szolgáltatásért, attól függően, hogy milyen személyes jellemzőkkel rendelkeznek, vagy mikor milyen időközönként, illetve hányadik alkalommal néztek meg egy adott terméket. A süti⁵⁷ segítségével a platform felismeri a vásárló digitális profilját, böngészési előzményeit, majd az erre létrehozott algoritmusok a profilra jellemző bizonyos sajátosságok alapján (impulzusvásárlásra való késztetések gyakorisága vagy például az IP cím alapján tartózkodási hely, lakhely, böngésző típusa, melyekből még az is kikövetkeztethető, hogy laptopról, vagy macbookról csatlakozik a felhasználó) az adott személyhez kalkulált személyre szabott árat mutatnak.⁵⁸ Ezeknél gyakoribb – és nehezebben szabályozható - eset a nem szándékos elfogultság, ami leggyakrabban a tanító adatkészletben már előzetesen is fellelhető előítéletek, torzítások, diszkriminatív trendek miatt alakul ki - minthogy ezek az adatok hűen tükrözik azon társadalmi és intézményi struktúrákat és normákat melyekben ilyen elfogultságok rejlenek. A harmadi kategória, esetén algoritmusok a tanulási folyamat során önmaguktól alakítanak ki elfogultságot, ami az adatokban rejlő mintázatok felismerésében, elsajátításán majd ezen felismerésekből levont következtetések döntési tényezőként történő elismerésén alapul.⁵⁹

3.3.3 Elfogult és diszkriminatív döntéshozatal a gyakorlatban

⁵⁷ A süti (cookie) egy olyan webszerver által létrehozott adatsomag, amit a felhasználó számítógépén tárolnak meghatározott ideig annak érdekében, hogy nyomon követhessék annak online tevékenységét, például azt, hogy honnan, milyen gyakran és milyen böngészővel látogat meg egy weboldalt. Az első webes sütit 1994-ben Lou Montulli, a Netscape Communications programozója alkotta meg, majd a 90-es évek második felében történő elterjedésüket követően hamar arra kezdték használni őket, hogy rögzítsék a mely oldalakat látogattak, milyen termékeket néztek meg vagy milyen hirdetésekre kattintottak a felhasználók. Később ahogy a profilozó algoritmusok fejlődtek a süti használata is egyre összetettebbé vált. A belőlük gyűjtött adatokból már új felhasználói interakciókra és online viselkedésmintákra is igyekeztek következtetéseket levonni. Lásd: Mayer, J., & Mitchell, J. C. (2012). Third-Party Web Tracking: Policy and Technology. IEEE Symposium on Security and Privacy

⁵⁸ E gyakorlatokról lásd bővebben: Bara Zoltán (2017): Dinamikus árazás az online kereskedelemben – Hogyan lehet hátrányos a fogyasztónak a dinamikus árazás? *Versenyügyek*. XIII. évfolyam (2), 4-19.o.

⁵⁹ Az algoritmikus elfogultságról, diszkriminációról lásd bővebben: Barocas, Solon – Selbst, Andrew D. (2016): Big Data's Disparate Impact. *California Law Review*. 104. 671–732.o. Azzal kapcsolatban, hogy az ilyen minták felismerése, majd szabályként elfogadása, hogy képes megerősíteni és reprodukálni a már létező faji, nemi, gazdasági egyenlőtlenségeket lásd bővebben: Eubanks, Virginia (2018): Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. New York, St Martin's Publishing. 14-39.o.

A döntéshozatalban megjelenő nem szándékos, rejtett torzításra⁶⁰ példaként szolgálhat az Amazon 2014-ben fejlesztett automatizált toborzó rendszere. Az algoritmust célja az volt, hogy előszűrje a meghirdetett technikai-mérnöki pozíciókra (pl. szoftverfejlesztői állás) a legalkalmasabb jelöltek listáját a nagy létszámú jelentkezői önéletrajzok közül. Hamar kiderült azonban, hogy a szelekció eljárása során az algoritmus elfogultságot mutatott a női jelentkezőkkel szemben. E probléma abból eredt, hogy az Amazon számítógépes modelljeit úgy képezték, hogy a jelentkezőket a céghez elmúlt tíz évben benyújtott önéletrajzok mintázatai alapján értékeljék. A rendszer azért mutatott elfogultságot a nőkkel szemben, mert a modell tanításához használt adatok a korábbi jelentkezők férfi túlsúlyát tükrözték. Ennek végeredményeként az algoritmus egyszerűen leminősítette azokat az önéletrajzokat, amelyek a "women's" (női) szót tartalmazzák, ahogy az történt például a "women's rugby team" (női rögbicsapat) esetében.⁶¹ Az Amazon később felhagyott az algoritmus használatával, elismerve, hogy nem tudták eredményesen kiküszöbölni a szoftver által visszatérő jelleggel tanúsított ilyen típusú elfogultságot.

Egy másik, gyakran idézett eset a diszkriminatív döntéshozatallal kapcsolatban az amerikai büntető igazságszolgáltatásban használt COMPAS rendszerhez kötődik.⁶² Az előre jelző algoritmus 137 gondosan összeválogatott kérdésre adott válasz, valamint az elkövető a bűnügyi nyilvántartásban tárolt adataiból származó információk kiértékelése alapján készített előrejelzést a vádlottak ismételt bűnelkövetésének valószínűségéről.⁶³ Az algoritmus alapú kockázatértékelés elsősorban abban segített a bírónak, hogy az adott ügyben döntsön az óvadék összegéről, illetve az előzetes letartóztatás vagy a feltételes szabadlábra helyezés megalapozottságáról. Egy ProPublica által készített elemzés arra a megállapításra jutott, hogy a COMPAS-rendszer értékelési eljárása során több esetben is elfogultságot mutatott az egyesült

⁶⁰ Az rejtett torzításnál (indirekt diszkriminációnál) az algoritmusok nem közvetlenül használják fel a védett tulajdonságokat (pl: nem, származás, vallás) hanem úgynevezett proxik (helyettesítő változók) segítségével következtetnek ezekre a jellemzőkre.

⁶¹ Jeffrey Dastin (2018): Insight - Amazon scraps secret AI recruiting tool that showed bias against women. REUTERS. Elérhető: <https://www.reuters.com/article/idUSKCN1MK0AG> (2021. 04. 11.)

⁶² A COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) rendszert a Northpointe fejlesztette ki és elsőként Wisconsin államban 2012-ben kezdték el használni a büntetőeljárások során. A programról bővebben: Park A. Lee (2019): Injustice Ex Machina: Predictive Algorithms in Criminal Sentencing. *UCLA Law Review*. Elérhető: <https://www.uclalawreview.org/injustice-ex-machina-predictive-algorithms-in-criminal-sentencing/> (2021. 07. 22.)

⁶³ Az algoritmus különböző szempontokból alapján elemzi az elkövetőt, figyelembe véve a korábbi bűncselekményeit, családi háttérét, társadalmi-gazdasági helyzetét, az életkörülményeinek stabilitását. Ezen információk alapján 1-től 10-ig terjedő skálán értékeli a visszaesés kockázatát, majd sorolja be a személyt magas, közepes vagy alacsony kockázatú kategóriába. Lásd: Gombos Katalin, Gyurancz Franciska Zsófia, Krausz Bernadett, Papp Dorottya (2021): A mesterséges intelligencia jogalkalmazási területen való hasznosíthatóságának alapjogi kérdései. In Török Bernát és Zódi Zsolt (szerk): A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 331-332.o.

államokbeli afroamerikai lakossággal szemben. Az adatok utólagos vizsgálatából tisztán kirajzolódott, hogy a program a nem visszaeső afroamerikai vádlottakhoz átlagosan közel kétszer magasabb kockázati százalékot rendelt (45%), mint a nem visszaeső fehér bőrű társaikhoz (23%). Emellett az elemzés arra is fényt derített, hogy a rendszer értékelése során azok a fehérbőrű vádlottak, akik újból bűncselekményt követtek el, közel kétszer olyan gyakran lettek alacsony kockázatúnak minősítve, mint az afroamerikai visszaesők (48% szemben 28%-kal).⁶⁴

Az egyenlő bánásmód megsértésének egy sajátos esete a bankok által előszeretettel használt hitelbírálati algoritmusokhoz kötődik.⁶⁵ Egy 2021-ben publikált tanulmány, arra a következtetésre jutott, hogy az ilyen kockázatértékelő rendszerek gyakran torzítanak a vékony hiteltörténettel (thin credit files) rendelkező ügyfelek hitelképességének megállapításakor. A vékony hiteltörténet kevés vagy korlátozottan rendelkezésre álló banki információra utal, ami gyakran előfordul azoknál az ügyfeleknél, akik nem igényeltek korábban kölcsönt, nem rendelkeznek hosszú hiteltörténettel, vagy kevés szolgáltatói terméket (pl. hitelkártyák) birtokolnak. A hitelminősítési algoritmusok nagymértékben támaszkodnak ezekre az információkra az ügyfél értékelésekor, ha valakinek a profiljához vékony előtörténet volt csatolva az megnehezítette a pontos kockázatbecslést, ami a gyakorlatban a lehetséges kockázatok túlbecslését, az ügyfél megbízhatóságának alul becslését eredményezte (bizonytalansági elfogultság).⁶⁶

⁶⁴ Az elemzés során több mint 10 000 bűnügyi vádlott esetét vizsgálták meg a floridai Broward megyéből, majd összevetették az algoritmus előrejelzéseit a ténylegesen bekövetkező eseményekkel. A tanulmány következtetései szerint megállapítható, hogy a COMPASS rendszer- az elfogultságra hajlamos döntéshozataltól eltekintve - az esetek többségében az afroamerikai és a fehér vádlottak esetében is körülbelül ugyanolyan arányban helyesen jósolta meg a visszaesést (59% a fehér vádlottaknál, 63% az afroamerikai vádlottaknál). Szignifikáns eltérés azoknál a fent is említett eseteknél volt megállapítható, ahol az előrejelzés téves következtetésre jutott. A ProPublica tanulmányának megállapításairól részletesebben: <https://www.propublica.org/article/how-we-analyzed-the-compas-recidivism-algorithm>

⁶⁵ A hitelminősítési algoritmusok a pénzügyi szektorban használt matematikai és statisztikai modellek, amelyek segítenek meghatározni az egyének vagy vállalatok hitelképességét és a hitelezéshez köthető lehetséges kockázatát. Az egyik legismertebb modell az USA-ban használatos FICO pontszám, mely 300-tól 850-ig terjed. Ezek az algoritmusok többféle adattípusra támaszkodnak, mint például az adós hiteltörténetére, jövedelmi adataira, adósságterheire és egyéb pénzügyi mutatókra, hogy felbecsüljék a kölcsönök időben történő visszafizetésének valószínűségét.

⁶⁶ A bizonytalansági elfogultság (uncertainty bias) okán a hitelbírálati algoritmusok hajlamosak a bizonytalanabb eseteket következetesen kockázatosabbnak értékelni. De nem csak az adatok mennyisége, hanem a minősége is meghatározó, így, ha valakinek hosszú, de ellentmondásos hiteltörténete van, az is növelheti a bizonytalanságot. A törzsszövegben említettekén túl ezzel kapcsolatban felmerül az a további probléma is, hogy ha egy adatbázisban bizonyos pozitív értékek alul reprezentáltak (pl. a hitelkérelemben olyan irányítószám szerepel, amellyel kevés megbízható banki ügyfél rendelkezik) akkor az algoritmus számára a róluk készült előrejelzések is automatikusan bizonytalanabbá válhatnak, mely eredményeképpen az algoritmus hajlamos lehet alacsonyabb összegű és rosszabb kondíciójú hiteleket felajánlani az ügyfélnek. A hitelkérelmek során felmerülő torzításokról és bizonytalansági elfogultságról bővebben: Banasik, J., Crook, J., & Thomas, L. (2003): Sample selection bias in credit scoring

Ezek az esetek is rámutatnak, hogy a megfelelő tanító adathalmaz létfontosságú az MI modellek számára, mivel azok minősége befolyásolja a modell pontosságát, megbízhatóságát és alkalmazhatóságát. Ha az MI elfogult, nem kellően reprezentatív adathalmazokon tanul, akkor hiányos, helytelen vagy torzított képet kap a valóságról.⁶⁷ A diszkrimináció tilalmát előíró hagyományos jogi keretek elsősorban a döntéshozatal végeredményére összpontosítanak. Azonban az MI esetében a folyamat minden lépése, az adatgyűjtéstől és elemzéstől kezdve a modellalkotásig és a tanításig, kritikus fontosságú a végső kimenet szempontjából. Elvárható tehát, hogy az MI-ra szabott jogi és etikai keretrendszereknek a döntést megelőző teljes folyamatra fókuszáljon, nem csupán magára a döntés eredményére. A szabályozó és ellenőrző szerveknek pedig olyan kérdésekre is összpontosítaniuk kell, mint az adatforrások megbízhatósága, az adatelemzési módszerek részrehajlása, valamint a folyamatos monitorozás és audit lehetőségének biztosítása.

3.3.4 Szabadság, demokrácia és a döntési autonómia

Az MI rendszerek elterjedése hamar ráirányította a figyelmet az emberi jogokkal és a demokratikus jogállamisággal kapcsolatos potenciális veszélyekre is. Ilyen kockázatok közé tartoznak többek között az egyenlő bánásmód követelményével, a tisztességes eljáráshoz való jog biztosításával, az adat- és magánélet védelmével, a modern megfigyelő állam kialakulásával, vagy az autonóm fegyverek megjelenésével kapcsolatos kihívások. Ezek kifejtésére ehelyütt nincs mód, azonban fontos szót ejteni egy olyan- a dolgozat későbbi részében is tárgyalandó – egyéni és társadalmi szinten is hosszútávú következményekkel bíró kockázatról, ami nagyban befolyásolhatja az egyének döntési autonómiájának és önrendelkezési jogának szabad gyakorlását.

models. *Journal of the Operational Research Society*, 54(8), 822–832.o. <https://doi.org/10.1057/palgrave.jors.2601578>. Az említett 2021-es tanulmány megállapításairól lásd bővebben: Blattner, L., és Nelson, S. (2021): How Costly is Noise? Data and Disparities in Consumer Credit. 27-30.o. ArXiv, abs/2105.07554.

⁶⁷ A torzított érzékelés közvetlen veszélyein túl fontos megemlíteni azt is, hogy az MI kimenetei gyakran olyan előrejelzések, amelyek közvetlen módon visszahatnak a környezetre, ezáltal hajlamosak fenntartani vagy még tovább mélyíteni a valóságban már meglévő egyenlőtlenségeket (un. önerősítő diszkriminatív hatást kifejtve) Például, ha egy prediktív rendőri rendszer az javasolja, hogy a rendőrség fokozottabban figyeljen meg egy számottevő etnikai kisebbséggel rendelkező városrészt. Ez több bűnözési adatot eredményezhet erről a területről és lakóiról. Az MI rendszer ezután még erősebb megfigyelést javasolhat erre a városrészt, mivel több "bűnözési" adatot gyűjt innen. Ezáltal a kisebbségek lakta terület fokozott monitorozása öngerjesztő, diszkriminatív kört hoz létre. Bővebben: Selbst, A. (2017): Disparate Impact in Big Data Policing. *Georgia Law Review*. 52, 3373.o. <https://doi.org/10.2139/SSRN.2819182>.

3.3.5 Kinek a választása?

Az önrendelkezés és a döntési autonómia olyan alapvető emberi jogok, amelyek az egyének szabadságát biztosítják saját életük tervezésében és döntéseik meghozatalában. Az MI azon képessége, mely az emberek preferenciáinak modellezésére és előrejelzésére irányul új kihívásokat jelent ezek vonatkozásában. Az elmúlt másfél évtizedben egyre több területen terjednek el algoritmikus ajánlórendszerek, amelyek kimondott célja, hogy gyorsabbá tegyék és megkönnyítsék az egyén döntéseit. Ezen eszközök használata azonban nagyban csökkentheti az egyén autonómiáját és kritikai gondolkodását, anélkül, hogy azt ő tolakodónak, vagy korlátozónak érezné (algorithmic micro-domination).⁶⁸ Az ajánlórendszerek működési elve, hogy hatalmas mennyiségű felhasználói adatot gyűjtenek és elemeznek, ⁶⁹ felismerve mintázatok és preferenciákat, melyek alapján személyre szabott javaslatokat tesznek az élet legkülönbözőbb területeken az útvonaltervezéstől kezdve, a diéták megalkotásán keresztül az online platformokon elérhető digitális tartalmak fogyasztásáig.⁷⁰ Első ránézésre ezek az ajánlások kényelmesebbé és hatékonyabbá tehetik a mindennapokat azáltal, hogy releváns opciókat kínálnak a végtelen elérhető kínálatból, azonban az algoritmusok által szűrt és prioritizált információk korlátozhatják az egyént abban, hogy az összes számára nyitva álló választási lehetőséget megismerje, szűkítve ezzel mozgásterét és látókörét.

⁶⁸ Danaher, J. (2019): The ethics of algorithmic outsourcing in everyday life. In K. Yeung & M. Lodge (szerk.) *Algorithmic Regulation* (91–118o.). Oxford University Press. 107.o. <https://doi.org/10.1093/oso/9780198838494.003.0005>

⁶⁹ Bár az algoritmikus ajánlórendszerek különböző adatokat és számolási modelleket használhatnak, ezek közül három terjedt el leginkább a gyakorlatban: (1) Tartalomszűrésen alapuló (content-based filtering) módszer, amely a felhasználó korábbi preferenciáiból és az ajánlásra szánt tételek tulajdonságaiból kiindulva tesz javaslatokat. Például egy zene-streaming szolgáltató platform, ahol a felhasználó korábban főleg klasszikus zenét hallgatott, erre alapozva fog hasonló stílusú, de a felhasználó számára még ismeretlen klasszikus zenei albumokat vagy előadókat ajánlani. Ez a szűrőmodell az tartalmat (mint szöveg, képek vagy hanganyag) elemezve keresi meg a hasonlóságokat, hogy olyan tételt ajánljon, amelyek illeszkednek a felhasználó által fogyasztott tartalmakhoz. (2) A demográfiai szűrés (demographic filtering) alapja az a feltételezés, hogy a hasonló személyes jellemzőkkel (mint nem, életkor, lakóhely) rendelkező felhasználóknak hasonlóak lehetnek az ízléseik és preferenciáik is. Ennek megfelelően a rendszer a felhasználók demográfiai hasonlóságai alapján tesz javaslatokat. Így például egy fiatal városi nő más ajánlásokat kap, mint egy idősebb vidéki férfi. (3) A kollaboratív szűrés (collaborative filtering) egy adott csoportnál megfigyelt preferenciákra támaszkodik az ajánlások generálásakor. A rendszer azonosítja azokat a felhasználókat, akik hasonló ízlésűek és preferenciákkal rendelkeznek, mint a "aktív" felhasználó, majd az ő értékeléseik és akcióik alapján tesz javaslatokat az aktív felhasználónak. Például, ha egy felhasználó 5 csillagosra értékelte Barcsi Tamás, című könyvét a bookline felületén, akkor a rendszer az adatfeldolgozás során azon vásárlókra összpontosít, akik hasonló értékelést adtak erre a könyvre, és az ő további olvasmányélményeik alapján fog újabb könyveket ajánlani a felhasználónak. Lásd: Bobadilla, J., Ortega, F. Hernando, A. and Gutierrez, A. (2013): *Recommender Systems Survey. Knowledge-Based System*. Vol(46), 112-113.o.

⁷⁰ Cheney-Lippold, John (2017): *We Are Data: Algorithms and the Making of Our Digital Selves*. New York: NYU Press. 88-92.o.

Az adatelemző technológiák fejlődése lehetővé tette, olyan részletes profilok létrehozását, melyek feltárják az adott egyénre jellemző kognitív sajátosságokat és döntéshozatali sebezhetőséget.⁷¹ Ezeket az információkat felhasználhatóak arra, hogy befolyásolják az egyének magatartását.⁷² Egy ilyen személyre szabott döntési környezetben a vásárlások ösztönzéstől kezdve a tudatos politikai mikrocélzásig számtalan lehetőség áll rendelkezésre az emberek manipulációjára.⁷³ E jelenségről a dolgozat második felében a GenMI rendszerek felemelkedése kapcsán még részletesebben is szó lesz.

Az egyéni autonómiával és döntési szabadságával összefüggésben, szintén kockázatot jelenthetnek, az állami- és magánszereplők által előszeretettel alkalmazott profilozáson, pontozáson és algoritmikus ajánlásokon alapuló társadalomszervező és döntéstámogató gyakorlatok. Ezek közül ehelyütt kettő kiemelése indokolt; (1) Hypernudge⁷⁴ (2) Általános társadalmi pontrendszerek.

(1): Richard Thaler és Cass Sunstein 2008-ban megjelent könyvükben a libertárius paternalizmus elméletével összhangban olyan ösztönzők (nudge) széles körű alkalmazása mellett érvelnek, amelyek az élet különböző területén segítik az embereket a helyes és kívánatos döntések meghozatalában.⁷⁵ Ezek az ösztönzők az emberi viselkedés előrejelezhetőségére és meghatározott kognitív torzításokra építve alakítanak ki olyan döntéskörnyezetet (choice environment), mely a kívánatos és előre megfogalmazott választás meghozásának kedvez. A döntéstervezés elméletét tovább gondoló, és azt az MI és Big Data által teremtett új társadalmi és technológiai környezetbe átültető Karen Yeung un. Hypernudge⁷⁶ koncepciójában felsejlik

⁷¹ A digitális profilok összetettségét jól szemlélteti, hogy már egy 2013-ban publikált tanulmány is arra következtetésre jutott, hogy a közösségi médián zajló interakciókból gyűjtött legalapvetőbb információk is képesek lehetnek az adott felhasználó személyes tulajdonságainak pontos meghatározására (pl.: életkor, nem, szexuális irányultság, vallási és politikai nézetek, személyiségjegyek, intelligencia, függőséget okozó szerek használata, szülőkkel való együttélés vagy különélés). Lásd: M. Kosinski, D. Stillwell, and T. Graepel (2013): Private traits and attributes are predictable from digital records of human behavior. *Proceedings of the National Academy of Sciences*, 110:5802–5805.; W. Youyou, M. Kosinski, and D. Stillwell (2015): Computer-based personality judgments are more accurate than those made by humans. *Proceedings of the National Academy of Sciences*, 112:1036–1040.o.

⁷² Harari Yuval Noah (2018): 21 Lecke a 21. századra. (ford: Torma Péter). Budapest, Animus kiadó. 55-60.o.

⁷³ Susser, D., Roessler, B., & Nissenbaum, H. (2019): Technology, autonomy, and manipulation. *Internet Policy Review*, 8(2), 1–21.o.

⁷⁴ Yeung, Karen (2017): Hypernudge: Big data as a mode of regulation by design. *Information, Communication and Society*, Vol.20(1)

⁷⁵ A szerzők a helyes döntés meghozatalát elsősorban a pénzügyek (pl.: nyugdíjtervezés, megtakarítások) és egészségügy területén támogatták, kiemelve, hogy az ösztönzés semmilyen esetben sem jelenthet egyet a választási szabadság korlátozásával, lehetőségek tiltásával. Bővebben: Thaler, R., Sunstein, C. (2008): *Nudge: Improving Decisions About Health, Wealth, and Happiness* London, Penguin Books

⁷⁶ Karen Yeung Hypernudge koncepciója nem csak ajánlásokat tesz, hanem aktívan alakítja és irányítja az egyének viselkedését, döntéseit. A rendszer három fő jellemzője, hogy (1) mindenütt jelen vannak, beágyazódva mindennapi környezetünkbe; (2) folyamatosan gyűjtenek adatokat az egyénről, hogy viselkedési modelleket

egy olyan jövőbeli automatizált magatartás-irányító algoritmikus rendszer képe, amelyben ezek az ösztönzők már tökéletesen kiismerik a célszemélyt, annak minden egyediségével és sebezhetőségével együtt. A döntéstámogató rendszerek automatizálása révén az egyén döntéshozatali környezete folyamatosan változna. Elképzelhető, hogy egy ilyen profilozáson, pontozáson, viselkedés alapú előrejelzésen alapuló rendszerben, ahol a köz és magánszolgáltatások elérésétől, a rendelkezésre álló bankhitelek és biztosítási konstrukciókon keresztül a fogyasztásra ajánlott termékekig bezárólag minden algoritmusok számítanak ki az egyén számára⁷⁷, ha egy választási lehetőséghez vagy szolgáltatáshoz való hozzáférés nem is nyílna meg azonnal, egy új, nagyobb pontossággal kiválasztott, és személyre szabottabb lehetőség kerülhet felkínálásra.⁷⁸ Ha elegendő felkínálandó lehetőség állna egy rendszer rendelkezésére, akkor a társadalom efféle értékelő számításokon alapuló irányítása egyáltalán nem látszana tényleges irányításnak, úgy tünne az egyénnek, hogy a szabadságába nem avatkozott be senki. Egy automatizált irányító társadalomban az ajánló gyakorlatok területei az élet minden területére kiterjedhetnének, és mivel a 21. század bővelkedik a felkínálható szolgáltatásokban, tartalmakban és termékekben, az egyén még a rideg számításokon alapuló, előre meghatározott döntések hálózatában sem érezheti autonómiájának vagy önrendelkezési jogának számottevő korlátozását.⁷⁹

(2) Az általános társadalmi pontrendszerek alkalmazásának legismertebb példája a kínai *Társadalmi Kreditrendszer*, melyben az automatikus értékelések egyenes következményeként bizonyos közszolgáltatásokhoz való hozzáférés is megtagadhatóvá válhat. A Kínai Államtanács által vázolt tervek alapján különböző minősítő és értékelő listák kialakítására kerül sor, melyben a jónak és kívánatosnak ítélt viselkedés – ilyen lehet például az önkéntes véradás – pontok hozzáadását, míg a nem megfelelő magatartás – például adófizetési határidő elmulasztása – pontlevonást vonhat maga után. Az összegyűjtött pontszámoktól függően részesülhet

építsenek; (3) ezek alapján személyre szabott beavatkozásokat hajtanak végre, hogy a kívánt irányba tereljék a célszemélyt. Ezáltal válnak képessé az összefüggő manipulációra (coherent manipulation). Yeung, Karen (2017): *Hypernudge: Big data as a mode of regulation by design. Information, Communication and Society*, Vol.20(1), 118-136.o.

⁷⁷ Az ehhez hasonló algoritmikus döntéseken alapuló közigazgatás- és társadalomszervezést nevezi Danaher algokráciának (algocracy). Szélsőséges esete lehet ha, az algoritmusok teljes mértékben automatizált döntéseket hoznak emberi felügyelet nélkül, ami potenciálisan kiszoríthatja az embereket a közügyek irányításából és a nyilvános döntéshozatali folyamatból, fokozatosan csökkentve ezzel a politikai folyamatok legitimációját és demokratikus jellegét is. Lásd bővebben: Danaher J. (2016): *The threat of algocracy: Reality, resistance and accommodation. Philosophy & technology* 29 (3), 245-268.o.

⁷⁸ Diósi Szabolcs (2022): *Behaviorally informed regulations- an emerging trend in modern public policymaking*. In Bendes Ákos L. et al. (szerk): *III. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai*. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar Doktori Iskola. 125-142.o.

⁷⁹ Barcsi Tamás - Diósi Szabolcs (2022): *Hasznos test, divídium, nyersanyag. A felügyeleti társadalomtól az (ön)ellenőrző és megfigyelési kapitalizmusig. Magyar filozófiai szemle* 66(3), 190-191.o.

jutalomban vagy büntetésben az egyén. Ilyen jutalom például az állam által támogatott tanfolyamokon való ingyenes részvétel, hozzáférés kedvezményes lakáshitelhez, olcsóbb tömegközlekedés biztosítása, míg a büntetésként az egyetemekre történő beiratkozás megtiltását, hitelkérelmek automatikus elutasítását, vagy a tömegközlekedés használatának határozatlan ideig történő megtiltását szabhatják ki.⁸⁰

Mindkét ismertetett példa esetében kiemelt kockázatot jelenthetne az, ha az ilyen rendszerek hatalomgyakorlás céljával megfelelő jogi szabályozás nélkül kerülnének alkalmazásra, hiszen ezek könnyen egész társadalmak kollektív magatartásbefolyásolásának eszközévé válhatnának.

3.4 Előre kalkulált kockázat, felülről szabályozott felhasználás

Érzékelve az MI elterjedésével összefüggő növekvő kockázatokat, az Európai Unió időben felismerte, hogy a technológiacsatlád hosszú távon is fenntartható és biztonságos alkalmazása érdekében mielőbb átfogó jogi szabályozásra van szükség. Az Unió 2017 óta élen jár az MI szabályozásával kapcsolatos törekvésekben, erőfeszítései e téren meghatározóak, hiszen ahogyan az a GDPR esetében is megfigyelhető volt, az Unió szabályalkotás hosszú távon is képes standardot szabni más szabályozás előtt álló országok számára (Brüsszeli hatás).⁸¹

Több éves jogalkotói eljárása során az Unió egy olyan jogi környezet kialakítására törekedett, amely nyitott a technológiai újítások felé, ösztönzi a kutatást és fejlesztést, egyúttal képes minimalizálni a társadalmi kockázatokat. Az innováció és európai versenyképesség megőrzésének támogatása, illetve az alapvető emberi jogok és európai értékek következetes, kompromisszumot nem ismerő védelme, mint olykor egymást korlátozó célok feloldására,

⁸⁰ A Kínai Államtanács 2014. június 14-én jelentette be, hogy elindítja nagyszabású, viselkedési adatokon, értékelő algoritmusokon és MI technológián alapuló Társadalmi Kreditrendszer programját. A pontrendszer a kínai bankok által évtizedek óta alkalmazott minősítési és rangsorolási rendszer masszív kiterjesztéseként működne, melynek gerincét a központi és regionális kormányzati szervek által, különböző szektorokban felhalmozott társadalmi/gazdasági magatartásokkal kapcsolatos adatok feldolgozása és azok meghatározott szempontok alapján történő értékelése adja. Erről részletesebben: Liu, Chuncheng (2019): Multiple Social Credit Systems in China. *Economic Sociology: The European Electronic Newsletter*, Vol 21 (1), 22–32.o.; Kollár Csaba (2020): Kína és a társadalmi kredit rendszere. *Hadtudomány: A magyar hadtudományi társaság folyóirata* 30:2, 79-97.o.; Vörös Zoltán (2019): Kína a digitális korszakban - Kínai internet és a Társadalmi Kreditrendszer. In: Dombi Judit – Rimai Dávid (szerk.): *Digitális forradalom világunkban: tanulmánykötet*. Pécs, Institutio Kiadó. 165-182.o.

⁸¹ A *Brüsszeli hatás* kifejezés arra utal, hogy az Európai Unió, mint jelentős és megkerülhetetlen gazdasági tényező - saját szabályozási törekvései révén - képes hosszútávon is irányadó globális szabványok kialakítására. Az elmúlt évtizedben ennek hatása tapasztalható volt többek között az adatvédelem, a digitális gazdaság, a fogyasztóvédelem és az élelmiszerbiztonság területén hozott szabályozásokban. Lásd bővebben: Bradford, Anu (2020): *The Brussels Effect: How the European Union Rules the World*. New York, Oxford University Press

pedig kockázatalapú szabályozási keretet javasolt, amely az MI rendszerek által jelentett kockázat mértékével arányos kötelezettségeket ír elő.

Az MI átfogó szabályozását nagyban nehezíti, hogy e rendszerek alkalmazási lehetőségének sokszínűsége miatt nehéz előre jelezni, hogy egy adott modell mely területeken okozhat fenyegetést a jövőben.⁸² Például az előző fejezetben említett magas minőségű képelemzésre képes mély neurális hálózat használható az orvostudományban rákdiagnosztikára, de szolgálhat a hatalomgyakorlás eszközeként is tömeges megfigyeléseken keresztül. A kockázatalapú megközelítés a kötelezettségek megállapításánál elsősorban arra koncentrál, hogy mely területeken alkalmazzák majd az adott MI-t.⁸³ Annak ellenére, hogy ez egy hatékony jogalkotói stratégia⁸⁴, GenMI felbukkanása új kihívásokat hozott e tekintetben is, mivel ezek a modellek már a felhasználói szinten is könnyen módosíthatóak, így lehetséges alkalmazási területüknek előre történő meghatározása – és az azzal kapcsolatos kockázatok kategorizálása – megközelítőleg lehetetlen feladat. Ennek orvoslására írtak hozzá – a jogalkotási folyamat kellős közepén - egy új fejezetet a törvényszöveghez és hoztak létre egy új MI kategóriát az uniós jogalkotók.

A következő fejezet részletesen tárgyalja az Európai Unió MI szabályozásával kapcsolatos jogalkotási törekvéseit. Mivel a jogalkotási folyamat párhuzamosan zajlott e dolgozat írásával, a fejezet tartalma is szervesen bővült egyidőben az uniós eljárás előrehaladásával. Ez lehetővé tette, a nyomon követését a kezdetektől eszközölt - hol kisebb hol nagyobb jogszövegváltoztatásoknak – valamint, hogy az azok mögött húzódó jogalkotási stratégiákat és elképzeléseket. A fejezet a szabályozási folyamat ismertetése során főként az alapvető emberi jogokra és európai értékekre veszélyes MI rendszerek és gyakorlatok ismertetésére és a releváns jogszabályok bemutatására összpontosít. Az MI-rendelet jövőbeli alkalmazásának és végrehajtásának technikai részleteit nem tárgyalja.

⁸² Bár egyes kockázatok utólag nyilvánvalóak lehetnek, azokat nem könnyű megvalósulásuk előtt azonosítani. Különösen nehéz előre látni az MI rendszerek jövőbeli alkalmazását (és az abból eredő kockázatokat), amikor azokat harmadik fél modelljeibe használják vagy integrálják.

⁸³ A megközelítés egy másik gyengesége, hogy mivel az MI rendszereket a szándékolt céljukat szem előtt tartva osztályozza, nem feltétlenül veszi figyelembe azokat a veszélyeket, amelyet egy MI rendszer a rendeltetési célján kívül történő használata okozhat. Ennek az a tovagyrűző hatása, hogy az e rendszerek fejlesztői és szolgáltatói viselik a szabályozási teher nagy részét. Ők felelősek a kockázatkezelési, adatminőségi, átláthatósági követelmények teljesítéséért. Ez jelentős jogi és adminisztratív terhet róhat rájuk, ami különösen a kisebb vállalatokra és startupokra terhes, amelyeknek eleve korlátozottabbak az erőforrásaik.

⁸⁴ Az EU újabban előszeretettel választja ezt a jogalkotási megközelítést a legújabb digitális technológiák szabályozására. Ez az MI szabályozás mellett az adatvédelem (GDPR) és az online platformok és közvetítő szolgáltatók esetében is megfigyelhető (DSA) volt.

4.1 Európa az MI szabályozás útján

Az Európai Tanács első alkalommal a 2017. október 19-ei ülésén elfogadott következtetésében kérte fel a Bizottságot arra, hogy terjesszen elő javaslatot „*a mesterséges intelligencia európai megközelítésére*” vonatkozóan.⁸⁵ A felkéréseknek eleget téve az Európai Bizottság 2018. április 25-én tette közzé első átfogó európai MI stratégiáját. E stratégia három fontos célt fogalmazott meg: (1) az EU technológiai és ipari kapacitásainak megerősítését; (2) a küszöbön álló társadalmi-gazdasági változásokra való felkészülést; illetve (3) az MI technológiákkal és gyakorlatokkal kapcsolatos megfelelő etikai és jogi keret kialakítását.⁸⁶ Az első kettő célt elsősorban az állami és magánbefektetések jelentős növelésével, kutatások és kutatóközpontok támogatásával, új képzések indításával, digitális készségek és tehetségek fejlesztésével, valamint az MI technológiák számára kedvező adatgazdag környezet (*adatökoszisztéma*) kialakításával tervezte elősegíteni. Az etikai és jogi keret tekintetében a Bizottság kiemelte, hogy az MI-nek mindenekelőtt az állampolgárok érdekeit kell szolgálnia, ezzel összefüggésben, fejlesztése és alkalmazása során mindenkor tiszteletben kell tartani mind az Európai Unióról szóló szerződés 2. cikkében meghatározott értékeket, mind az EU Alapjogi Chartájában felsorolt alapvető jogokat.⁸⁷ A Bizottság továbbá hangsúlyozta, hogy az MI alkalmazásoknak nemcsak a jogszabályoknak kell megfelelniük, hanem bizonyos etikai elvekkel is összhangban kell lenniük. Ezen célok előmozdítása érdekében a Bizottság 2018-ban létrehozta a *Mesterséges intelligenciával foglalkozó magas szintű szakértői csoportot*, amit az etikai iránymutatások összeállításával, illetőleg átfogóbb MI-politikára vonatkozó ajánlások kidolgozásával bízott meg.⁸⁸

⁸⁵ Európai Tanács (2017): Következtetések, EUCO 14/17, 7.o.

⁸⁶ Európai Bizottság (2018): Mesterséges intelligencia Európa számára (A közös európai adattér felé) COM (2018) 237 final. 7-20.o.

⁸⁷ A Bizottság egy 2019-ben közzétett kommentárjában ezt a megközelítést nevezi "emberközpontú" mesterséges intelligenciának. Ennek értelmében az MI *nem maga a cél, hanem egy, az embereket szolgáló eszköz, melynek végső célja az emberi jólét növelése*. Erről részletesebben: Európai Bizottság (2019): Az emberközpontú mesterséges intelligencia iránti bizalom növelése COM(2019) 168 final.

⁸⁸ A szakértői csoport 2019 márciusában adta át a Bizottságnak az iránymutatásuk végleges változatát, melyben a megbízható MI kialakítását három előfeltételhez szabta: (1) a mesterséges intelligenciának be kell tartania a jogszabályokat; (2) meg kell felelnie az etikai elveknek; (3) stabilnak kell lennie. Ezt figyelembe véve a szakértői csoport hét olyan kulcsfontosságú követelményt sorolt fel, melyeknek az MI-rendszer teljes életciklusa során kötelező megfelelni: (1) az emberi cselekvőképesség támogatása és emberi felügyelet; (2) műszaki stabilitás és biztonság; (3) adatvédelem és adatkezelés; (4) átláthatóság; (5) sokféleség, megkülönböztetésmentesség és

Az áprilisban ismertetett stratégiára építve az Európai Tanács a 2018. június 28-i ülésén elfogadott következtetésben⁸⁹ felkérte a Bizottságot, hogy a tagállamokkal együttműködve dolgozzon ki egy *összehangolt MI tervet*. A Bizottság ugyanezen év decemberében mutatta be a *Mesterséges intelligenciára vonatkozó összehangolt tervet*⁹⁰, melyben hangsúlyozta, hogy a vázolt célok elérése érdekében – például a beruházások növelése, kiterjedt adat-ökoszisztéma létrehozása, az EU digitális versenyképességének hosszú távú biztosítása – európai szintű koordinációra, határokon átnyúló szorosabb együttműködés kiépítésére és nemzeti jogalkotási töredezettségektől mentes egységes piac kiépítésére van szükség. A közös együttműködés alapjainak megteremtésén és a legfontosabb beruházási területek meghatározásán túl a Bizottság arra is ösztönözte a tagállamokat, hogy dolgozzanak ki nemzeti stratégiai elképzeléseket az MI-ről.

Az MI szabályozással kapcsolatos uniós törekvések 2019-ben kaptak nagyobb nyilvánosságot amikor Ursula von der Leyen frissen megválasztott Bizottsági elnök a 2019–2024 közötti időszakra szóló, *Ambiciózusabb Unió*⁹¹ című politikai iránymutatásában kijelentette, hogy hivatali idejének *első száz napjában* a Bizottság jogszabályt fog előterjeszteni az MI emberi és etikai vonatkozásaival kapcsolatos összehangolt európai megközelítésre vonatkozóan.

A Bizottság végül 2020. február 19-én⁹² tette közzé az MI-vel kapcsolatos *Fehér Könyvét*, amelyben *konkrét szakpolitikai alternatívákat*⁹³ fogalmazott meg arra vonatkozóan, hogy miként lehet összeegyeztetni a magas minőségű MI technológiák és gyakorlatok fokozottabb

méltányosság; (6) társadalmi és környezeti jólét; (7) elszámoltathatóság. Lásd bővebben: Európai Bizottság, Mesterséges intelligenciával foglalkozó magas szintű szakértői csoport (2019): Etikai iránymutatás a megbízható mesterséges intelligenciára vonatkozóan, 19-24.o.

⁸⁹ Európai Tanács, Az Európai Tanács ülése (2018. június 28.) – Következtetések, EUCO 9/18 2018, 8.o.

⁹⁰ Európai Bizottság (2018): A mesterséges intelligenciáról szóló összehangolt terv COM(2018) 795 final.

⁹¹ *Hivatali időm első száz napjában a mesterséges intelligencia humán és etikai következményeire vonatkozó összehangolt európai megközelítésről szóló jogszabályt fogok előterjeszteni. Ez arra is ki fog terjedni, hogy hogyan tudjuk felhasználni az óriási méretű adathalmazokat, amelyek gazdagságot teremtenek a társadalmaink és vállalkozásaink számára.* Lásd bővebben: Európai Bizottság, Kommunikációs Főigazgatóság, Leyen, U. (2019): *Ambiciózusabb Unió – Programom Európa számára: politikai iránymutatás a hivatalba lépő következő Európai Bizottság számára (2019–2024).* Elérhető: <https://data.europa.eu/doi/10.2775/56352>

⁹² A *Fehér Könyv*vel egy időben a Bizottság közzétette a "Jelentés a mesterséges intelligencia, a tárgyak internete és a robotika biztonsági és felelősségi vonatkozásairól" című dokumentumot, valamint az Európa digitális jövőjének alakításáról (COM(2020) 0067) és az Európai Adatstratégiáról (COM(2020) 0066) szóló közleményeit is.

⁹³ A *Fehér Könyv* a Bizottság által 2018 áprilisiában közzétett MI stratégiájára építve tovább hangsúlyozza a beruházások növelésének (magánszektor bevonása, kkv-k fokozatos támogatása, közigazgatási szervek technológiai átalítása) és az adatgazdag környezet kialakításának szükségességét. A koordinált terv alapján pedig az egységes piac létrehozására és szétterjedt nemzeti jogalkotás elkerülésének szükségességére helyezi a hangsúlyt. A szakértői tanács által készített etikai iránymutatást következtetéseit elfogadva külön kiemeli, hogy az MI technológiákkal kapcsolatos etikai, jogi problémákból az átláthatatlanság („feketedoboz-hatás”), az összetettség, a kiszámíthatatlanság az autonóm viselkedés, az elfogultságot és megkülönböztetést követő gyakorlatok, illetve a magánszféra és a személyes adatok védelme kitüntetett figyelmet igényelnek.

elterjedését (*kiválósági ökoszisztéma*) a biztonságos alkalmazási környezet kiépítésével (*bizalmi ökoszisztéma*).⁹⁴ A Bizottság hangsúlyozta, hogy az MI új szabályozási keretének kialakításánál fontos követelmény a megfelelő mértékű szigor biztosítása, ugyanakkor lényeges szempont a joganyag túlzottan előíró jellegűvé válásának elkerülése is, amely az aránytalanul nagy terhek megállapításával hosszútávon az innováció gátjaként szolgálna. A hatékony de arányos beavatkozás lehetőségének megteremtése érdekében a *Fehér Könyv* bevezette a *magas kockázatú MI-rendszerek* fogalmát⁹⁵ melyekkel szemben szigorúbb követelményeket állított (többek között stabilitással és pontossággal, átláthatósággal, tájékoztatással, az emberi felügyelettel és az adatkormányzással kapcsolatosan). A szabályozási javaslatok ismertetésén túl a *Fehér Könyv* közzétételével egy időben a Bizottság széles körű nyilvános konzultációt⁹⁶ is indított az MI európai megközelítésére vonatkozó szakpolitikai és szabályozási intézkedésekről.

A konzultáción elhangzottakra, illetve az Uniós intézmények által korábban közzétett állásfoglalásokra, javaslatokra, iránymutatásokra és stratégiákra építve a Bizottság végül 2021. április 21-én tette közzé az MI szabályozásával kapcsolatos rendelet-tervezetét.

4.2 A kockázat alapú megközelítés

Összegezve a közzétételt megelőző években megfogalmazott prioritásokat, a rendelet-tervezet a következő négy általános célkitűzést fogalmazta meg: (1) „*annak biztosítása, hogy az Unióban forgalomba hozott és használt MI-rendszerek biztonságosak legyenek, és tiszteletben*

⁹⁴ Európai Bizottság: Fehér könyv a mesterséges intelligenciáról – A kiválóság és a bizalom európai megközelítése COM(2020) 65 final. 3-4.o.

⁹⁵ A *Fehér Könyv* javaslata alapján egy MI-rendszert abban az esetben kell nagy kockázatúnak tekinteni, ha; (1) olyan ágazatban használják, ahol jelentős kockázatok várhatók (pl.: energiaszektor, határellenőrzés, igazságszolgáltatás); (2) az adott ágazatban ezenfelül olyan módon is használják az adott alkalmazást, ami jelentős kockázatokat hordoz. Lásd bővebben: Európai Bizottság (2020): Fehér könyv a mesterséges intelligenciáról: a kiválóság és a bizalom európai megközelítése. COM(2020) 65 final, 21-29.o.

⁹⁶ A nyilvános konzultációhoz, melynek lefolytatására online keretek között került sor 2020. február 19. és június 14. között, összesen 1.215 résztvevő (magánszemélyek, vállalatok, kutatói csoportok, hatóságok) csatlakozott. Összességében elmondható, hogy a válaszadók többsége kifejezetten támogatta a kockázatalapú megközelítést. Az ilyen megközelítés alkalmazását jobb megoldásnak tekintették, mint a valamennyi MI-rendszert érintő általános szabályozást. A résztvevők emellett válaszaikban - elfogadva a *Fehér Könyv* és a szakértői csoport által megfogalmazott legfőbb veszélyforrásokat - további konkrét kockázatokkal kapcsolatban is hangot adtak aggodalmaiknak, például az online hirdetések során tapasztalt differenciált ár képzés, a kommunikációs szűrőbuborékok kialakulása, a profilalkotáson alapuló manipuláció kockázata, a politikai folyamatokba és választásokba való beavatkozás veszélye, illetve az automatizált döntésekhez kapcsolódó állampolgári kontroll elvesztése kapcsán.

tartsák az alapvető jogokra és az uniós értékekre vonatkozó hatályos jogszabályokat”; (2) „a jogbiztonság biztosítása a mesterséges intelligenciába történő beruházások és a mesterséges intelligenciát érintő innováció elősegítése érdekében”; (3) „az irányításnak és az MI-rendszerek tekintetében az alapvető jogokra és biztonsági követelményekre vonatkozó hatályos jogszabályok hatékony érvényesítésének a javítása”; (4) „a jogszerű, biztonságos és megbízható MI-alkalmazások tekintetében az egységes piac kialakításának elősegítése és a piac széttöredezettségének megelőzése”.⁹⁷

Az olykor eltérő stratégiákat igénylő célkitűzések (pl.: innováció, fejlesztés és a beruházások támogatása vs. az uniós értékek valamint az állampolgárok szabadságjogainak kitüntetett védelme) sikeres elősegítése érdekében a rendelet-tervezett kockázatalapú megközelítést alkalmazott az MI-rendszerek Unión belül történő fejlesztésének, forgalmazásának és használatának szabályozására. E megközelítés elődleges célja, hogy időben azonosítsa és értékelje az MI-rendszerekhez kapcsolódó lehetséges kockázatokat, valamint, hogy a potenciálisan felmerülő kockázatok esetén olyan *minimumkövetelményekre* korlátozódjon, amelyek képesek hatékonyan kezelni a veszélyt, de egyúttal nem vezetnek *szükségtelen kereskedelmi korlátozáshoz*. Az arányossági alapokon nyugvó felfogás értelmében a javaslat az MI által jelentett potenciális kockázat növekedésével egyre szigorúbb jogszabályi követelmények alkalmazását írta elő. Ennek megfelelően tiltani rendelte az MI egyes különösen – *elfogadhatatlan mértékben* – veszélyes felhasználási módjait, szigorúan szabályozta a nagy kockázatú MI rendszerek fejlesztését, forgalmazását, használatát, illetve enyhe követelményeket állapított meg a kevésbé kockázatos rendszerek számára. A rendelet-tervezet négy kockázati szintet határozott meg, amelyek az (1) elfogadhatatlan, (2) nagy, (3) korlátozott és a (4) minimális kockázatú MI rendszereket jelölik.

4.3 Korlátozott és minimális kockázatú MI-rendszerek

Mínthogy alkalmazásuk nem jelent kirívó kockázatot az állampolgárok alapvető jogaira és biztonságára nézve, a rendelet-tervezet a korlátozott és minimális kockázatú MI-rendszerekkel kapcsolatosan enyhe jogkövetelményekre korlátozódott.

A minimális kockázatú rendszerek (például az e-mailszolgáltatók által alkalmazott spam-szűrők) Európai Unión belül történő fejlesztése és használata a hatályos jogszabályok alapján, további jogkötelezettségek nélkül engedélyezett. A Bizottság e rendszerekkel kapcsolatban

⁹⁷ MI rendelet-tervezet (Indoklás 1.1.)

elsősorban *ösztönözte* az olyan magatartási kódexek *önkéntes* kidolgozását és alkalmazását, amelyek célja, hogy előmozdítsák a veszélyesebb MI-rendszerekkel kapcsolatosan megfogalmazott követelmények betartását.⁹⁸

A rendelet-tervezet 52. cikkében foglalt átláthatósági kötelezettség azonban a *természetes személyekkel* közvetlen módon érintkező, nem nagy kockázati besorolású MI-rendszerek tekintetében is kötelezővé tette a szolgáltatók számára annak biztosítását, hogy „*az MI-rendszereket úgy tervezzék meg és fejlesszék ki, hogy a természetes személyek tájékoztatást kapjanak arról, hogy valamely MI-rendszerrel állnak kapcsolatban, kivéve, ha ez a körülmények és a használat kontextusa alapján nyilvánvaló*”. Az 52. cikk továbbá rendelkezett arról is, hogy azon MI-rendszerek felhasználóinak, melyek képesek valós személyekre, helyekre vagy eseményekre megtévesztő módon hasonlító *képi, audio-, illetve videótartalmat generálni vagy manipulálni*, kötelezően fel kell tüntetniük, hogy az általuk közölt tartalmak mesterségesen lettek létrehozva.⁹⁹ Ezek az átláthatósági szabályok a minimális kockázatú MI-alkalmazások kapcsán elsősorban a különböző digitális/virtuális asszisztensek, chatbotok, illetve a későbbiekben részletesen is kifejtésre kerülő *deepfake* tartalmak előállítására alkalmas rendszerek potenciális veszélyeinek mérséklésére szolgáltak. Ezekről a technológiákról a későbbiekben még részletesen is szó lesz.

4.4 Nagy kockázatú MI-rendszerek

A Bizottság a rendelet-tervezet Indoklásában előzetesen rögzítette, hogy azon MI-rendszerek tartoznak ebbe a kategóriába, *amelyek nagy kockázatot jelentenek a természetes személyek egészségére és biztonságára vagy alapvető jogaira nézve*.¹⁰⁰ Ezt követően részletesebben a III. Címben foglalkozik a magas kockázatú MI-rendszerekkel, melyeknek két fő kategóriáját határozza meg: (1) azon rendszerek, melyeket *már létező uniós harmonizációs jogszabályok hatálya alá tartozó termék biztonsági alkatrészeként vagy önmagában ilyen termékként kívánják használni*,¹⁰¹ (2) azon MI-rendszerek, melyek felhasználási módjukból vagy

⁹⁸ MI rendelet-tervezet 69. cikk

⁹⁹ MI rendelet-tervezet 52. cikk

¹⁰⁰ MI rendelet-tervezet (Indoklás 5.2.3.)

¹⁰¹ A rendelettervezet III. Címének 6. cikke tartalmazza a nagy kockázatú MI-rendszerekre vonatkozó besorolási szabályokat. Ennek értelmében az MI-rendszer nagy kockázatúnak minősül, ha mindkét alábbi feltétel teljesül: (1) a javaslat II. mellékletben felsorolt uniós ágazati jogszabályok *hatálya alá tartozó termék biztonsági alkatrészeként vagy önmagában ilyen termékként kívánják használni*; (2) a javaslat II. mellékletében felsorolt uniós

alkalmazási területükből kifolyólag számottevő módon befolyásolhatják az állampolgárok alapvető szabadságjogait. A Bizottság a javaslatához csatolt III. mellékletben nevesítette e nyolc területet: (1) természetes személyek biometrikus azonosítása és kategorizálása; (2) kritikus infrastruktúra irányítása és működtetése; (3) oktatás és szakképzés; (4) foglalkoztatás; (5) alapvető magánszolgáltatásokhoz és közszolgáltatásokhoz való hozzáférésről való döntés, (6) bűnüldözés (7) migráció, menekültügy és határellenőrzés; (8) igazságszolgáltatás és demokratikus folyamatok.¹⁰² A rendelet időtállóságának és jövőbeni alkalmazhatóságának elősegítése érdekében a javaslattervezet 7. cikke arra is felhatalmazza a Bizottságot, hogy a III. mellékletben nevesített MI-rendszerek jegyzékét - ilyen például a bűnüldöző hatóságok által a bűnisméltés valószínűségének megállapítására, a potenciális bűncselekmények előrejelzésére, vagy a természetes személyek érzelmi állapotának észlelésére használt MI-rendszerek listája - a jövőben további MI-k hozzáadásával is bővítse, amennyiben azokat az előbb felsorolt nyolc terület valamelyikén kívánják használni, illetve amennyiben azok alkalmazása az *egészségben és a biztonságban okozott kár vagy az alapvető jogokat érő kedvezőtlen hatás tekintetében* kiemelt kockázattal járnak.¹⁰³¹⁰⁴

A rendelet-tervezet III. címének 2. fejezete rögzítette a magas kockázatú MI-rendszerekre vonatkozó jogszabályi követelményeket. A javaslat többek között kötelezettségeket fogalmazott meg a *kockázatkezelési rendszerek* létrehozásával, magas minőségű

harmonizációs szabályok értelmében *forgalomba hozatala vagy üzembe helyezése céljából harmadik fél által végzett megfelelőségértékelésnek kell alávetni.*

¹⁰² A III. melléklet részletesebben is nevesít alkalmazási területeket és módozatokat a nagy kockázatú MI rendszerekkel kapcsolatban, ilyenek például: (1) a természetes személyek valós idejű és nem valós idejű távoli biometrikus azonosítására szolgáló MI-rendszerek; (2) a közúti forgalom irányításában és működtetésében használt MI-rendszerek; (3) oktatási és szakképzési intézmények hallgatóinak értékelésére, vagy az oktatási intézményekbe való felvételhez szükséges tesztek résztvevőinek értékelésére szolgáló rendszerek; (4) természetes személyek toborzására, kiválasztására, üres álláshelyek meghirdetésére, pályázatok szűrésére, valamint a jelöltek interjúk során történő értékelésére szolgáló MI-rendszerek; (5) állami segítségnyújtási ellátásokra és szolgáltatásokra való jogosultságának értékelésére, vagy a természetes személyek hitelképességének értékelésére, hitelpontszámuk megállapítására szolgáló MI-rendszerek; (6) a bűnüldöző hatóságok által a természetes személyek bűnelkövetésének kockázatértékelésére használt rendszerek, vagy olyan rendszerek melyeket profilalkotás alapján természetes személyek személyiségbeli jellemzőinek és tulajdonságainak vagy múltbeli bűnöző magatartásának értékelése alapján bűncselekmények előfordulásának vagy megismétlődésének előrejelzésére használnak; (7) határellenőrzés során valamely tagállam területére belépett vagy oda belépni szándékozó természetes személy által jelentett kockázat értékelésére használt rendszerek; (8) olyan MI-rendszerek, amelyek célja, hogy segítsék az igazságügyi hatóságokat a tények és a jog kutatásában és értelmezésében, valamint a jog konkrét tényállásra történő alkalmazásában. MI rendelet-tervezet III. Melléklet 1-8.

¹⁰³ MI rendelet-tervezet 7. cikk

¹⁰⁴ Többben fogalmaztak meg kritikát a tervezetben biztosított bővítési eljárással kapcsolatban. Smuha és társai például a mechanizmus anti-demokratikus jellegét kifogásolva, azt javasolták, hogy a jogszabály későbbi változtatában fontos volna konzultációs és részvételi jogokat is biztosítani az EU minden állampolgára számára a magas kockázatú rendszerek listájának módosítását illetően. Lásd bővebben: Smuha, Nathalie A., Ahmed-Rengers, Emma Harkens, Adam, Li, Wenlong MacLaren, James Piselli, Riccardo, Yeung, Karen (2021): How the EU Can Achieve Legally Trustworthy AI: A Response to the European Commission's Proposal for an Artificial Intelligence Act. <http://dx.doi.org/10.2139/ssrn.3899991>

*adatkormányzással és adatgazdálkodással, átlátható működéssel, a használat teljes időtartalma alatt fenntartott hatékony emberi felügyelettel, valamint a pontos és biztonságos tervezéssel és fejlesztéssel kapcsolatban.*¹⁰⁵ Ezen felül a rendelet-tervezet a szolgáltatókra és felhasználókra tekintettel is írt elő további kötelezettségeket, ilyen például a *minőségirányítási rendszer* bevezetése vagy a kötelező *megfelelőségértékelés eljárás* lefolytatása az MI-rendszer forgalomba hozatalát illetve üzembe helyezését megelőzően.¹⁰⁶ A rendelet-tervezet 71. cikke értelmében abban az esetben ha az adott MI-rendszer nem felel meg a felsorolt követelményeknek vagy kötelezettségeknek, *legfeljebb 20 000 000 EUR összegű közigazgatási bírsággal, illetve vállalkozások esetében az előző pénzügyi év teljes éves világpiaci forgalmának legfeljebb 4 %-át kitevő összeggel sújtható* (a magas minőségű adatkormányzás és adatgazdálkodási követelményének megszegése esetén pedig 30 000 000 EUR összegű közigazgatási bírsággal, vagy a teljes éves világpiaci forgalmának legfeljebb 6 %-át kitevő összeggel).¹⁰⁷

4.5 Tiltott gyakorlatok, „elfogadhatatlan kockázat”

A Bizottság álláspontja szerint bizonyos MI-rendszerek olyan mértékű – *elfogadhatatlan* - kockázatot jelentenek az uniós értékek és állampolgárok alapvető jogaira nézve, melyeknek orvoslása nem lehetséges sem technikai megoldásokkal, sem különböző jogi és eljárási biztosítékokkal; az uniós állampolgárok érdekeinek védelmében ezen alkalmazások teljes tiltása szükséges. A rendelet-tervezet négy tiltott kategóriát nevez meg, amelyek közül hármat teljes egészében tiltani rendel, egyet pedig csak meghatározott feltételek teljesülése esetén engedélyez.

4.5.1 Magatartást torzító gyakorlatok

Az V. Címben megnevezett négy tilalom közül az első kettő a különböző emberi magatartásokat manipuláló vagy azokat befolyásoló MI alkalmazásokra irányult. Ennek értelmében tilos az olyan rendszerek forgalomba hozatala, üzembe helyezése vagy használata, amelyek:

¹⁰⁵ MI rendelet-tervezet III. Cím 2. fejezet, 8-15. cikk

¹⁰⁶ MI rendelet-tervezet III. Cím 3. fejezet, 16-29. cikk

¹⁰⁷ MI rendelet-tervezet X. Cím 71. cikk

„(a) szubliminális technikákat alkalmaznak az adott személy tudatán kívül, annak érdekében, hogy lényegesen torzítsák egy személy magatartását oly módon, amely e személynek vagy más személynek testi vagy lelki károsodást okoz vagy okozhat;

(b) a személyek egy meghatározott csoportjának életkor, testi vagy szellemi fogyatékoság miatt fennálló valamilyen sebezhetőségét kihasználják annak érdekében, hogy torzítsák az adott csoporthoz tartozó személy magatartását oly módon, amely e személynek vagy más személynek testi vagy lelki károsodást okoz vagy okozhat”¹⁰⁸

A Bizottság által közzétett javaslatban a magatartást befolyásolni és torzítani képes MI-rendszerek forgalomba hozatala, üzembe helyezése vagy használata csak abban az esetben volt tiltott, amennyiben azok az egyénnek károsodást okoznának. Az (a) pont esetében például nem a befolyásolás ténye vagy az arra irányuló szándék volt a mérvadó, hanem az, hogy a gyakorlat végrehajtása olyan módon történjen, „amely e személynek vagy más személynek testi vagy lelki károsodást okoz vagy okozhat”. A javaslatszövegbe beemelt szigorú testi és lelki károkra való korlátozás számos kiskaput rejtett magában, mely hosszú távon alááshatta volna a rendelet eredeti célkitűzéseit. Ahogy Veale is utalt rá, e megszorító hatály ellen joggal volt felhozható, hogy az állampolgárokat védelmezni szükséges valamennyi olyan szubliminális (tudatalatti) technikával vagy befolyással szemben, amely megkerüli, gyengíti vagy aláássa a racionális kontrollt, függetlenül attól, hogy ezek a befolyásolási kísérletek tényleges testi vagy lelki károkat okoznak vagy sem.¹⁰⁹ Ugyanezen gondolat mentén a (b) pont tekintetében megalapozott lehet azon elvárás is, hogy amennyiben egy MI-rendszert arra terveztek vagy azt oly módon alkalmazzák, hogy kihasználja bizonyos sérülékeny csoportok sebezhetőségét (gyermekek, idősek, értelmi vagy testi fogyatékosággal élő emberek) akkor azt – tekintet nélkül annak tényleges következményeire – be kell tiltani.

A Régiók Európai Bizottsága egy 2021 decemberében közzétett véleményében javasolta a kár feltétel kiterjesztését a testi és lelki károkon túl a gazdasági károkra is.¹¹⁰ Így a tilalom védelmet nyújthat azon MI technikákkal szemben, amelyek a kiszolgáltatott emberek számára negatív gazdasági következményekkel, *pénzügyi veszteséggel vagy gazdasági diszkriminációval*

¹⁰⁸ MI rendelet-tervezet II. Cím 5. cikk (1) a-b

¹⁰⁹ Veale, Michael - Frederik Zuiderveen Borgesius (2021): Demystifying the Draft EU Artificial Intelligence Act. Computer Law Review International, 22(4) 97-112.o.

¹¹⁰ Az Régiók Európai Bizottsága az alábbi kiegészítést javasolta az első bekezdés (a) pontja tekintetében: „testi vagy lelki károsodást okoz vagy okozhat, sérti vagy sértheti más személynek vagy személyek egy csoportjának alapvető jogait, beleértve testi vagy lelki egészségüket és biztonságukat, károkat okoz vagy okozhat a fogyasztóknak, például pénzügyi veszteséget vagy gazdasági diszkriminációt, vagy aláássa vagy alááshatja a demokráciát és a jogállamiságot” Régiók Európai Bizottsága (2021): Vélemény. A mesterséges intelligenciával kapcsolatos európai megközelítés – A mesterséges intelligenciáról szóló jogszabály. SEDEC-VII/022, 13.o.

járhatnak (pl. az arra hajlamos egyént játékfüggőségbe sodorja). A (b) pont tekintetében szintén indokolt lenne kiterjeszteni a védett személyek körét mely jelenlegi formájában a *személyek egy meghatározott csoportjának életkor, testi vagy szellemi fogyatékoság miatt fennálló valamilyen sebezhetőségére* korlátozódik. A rendelet-tervezetben felsorolt személyek körén kívül sebezhetőek – egyúttal könnyebben befolyásolhatóak és kihasználhatóak – lehetnek azon személyek is, akik különféle mentális betegségekkel vagy viselkedési zavarokkal küzdenek (pl.: depresszió, bipoláris zavar, generalizált szorongás, ADHD vagy függőség).

Az Európai Gazdasági és Szociális Bizottság pedig azt vetette fel, hogy a szubliminális technikák nemcsak testi és lelki károkhoz vezethetnek, hanem tágabb kontextusban nézve *„alkalmazási környezetüktől függően egyéb káros személyi, társadalmi vagy demokratikus hatásokkal – például a szavazási magatartás megváltozásával – is járhatnak”*. Ennek kapcsán javasolta az EGSZB a 2021 szeptemberében elfogadott véleményében, hogy a tiltott gyakorlatok körét oly módon bővítsék ki, hogy azok magukba foglalják azon alkalmazásokat is, melyek egy adott személynek vagy személyek csoportjának *„sérti vagy sértheti alapvető jogait, illetve ezen túl sérti vagy sértheti az általánosabb/össztársadalmi szinten értelmezhető demokráciát és a jogállamiságot”*.¹¹¹

4.5.2 Általános társadalmi pontozás

A tilalmak második csoportját az úgynevezett *általános társadalmi pontozással* kapcsolatos alkalmazások alkotják. Ennek értelmében tilos az:

*„MI-rendszerek hatóságok általi vagy nevükben történő forgalomba hozatala, üzembe helyezése vagy használata természetes személyek megbízhatóságának értékelésére vagy osztályozására egy bizonyos időszakon keresztül, közösségi magatartásuk, illetve ismert vagy előre jelzett személyes vagy személyiségi tulajdonságok alapján ...”*¹¹²

¹¹¹ Európai Gazdasági és Szociális Bizottság (2021): Vélemény - Javaslat európai parlamenti és tanácsi rendeletre a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról INT/940, 4-5.o

¹¹² *„oly módon, hogy a közösségi pontszám a következő helyzetek egyikéhez vagy mindkettőhöz vezet: (i) a gyakorlat természetes személyek vagy azok egész csoportjával szemben hátrányos vagy kedvezőtlen bánásmóddhoz vezet olyan szociális kontextusban, amely nem függ össze azzal a kontextussal, amelyben az adatokat eredetileg létrehozták vagy gyűjtötték; (II) a gyakorlat természetes személyek vagy azok egész csoportjával szemben olyan hátrányos vagy kedvezőtlen bánásmóddhoz vezet, amely indokolatlan vagy aránytalan közösségi magatartásukhoz vagy annak súlyosságához képest”* MI rendelet-tervezet II. Cím 5. cikk i,ii

A Bizottság hamar felismerte a pontozáson alapuló gyakorlatokban rejlő jelentős társadalmi veszélyeket. Az ehhez hasonló *általános társadalmi pontozási* gyakorlatok korai és szigorú szabályozásának fontosságát – minthogy azok az állami túlhatalom kiépítésének kockázatát rejtik - az Európai Unió több alkalommal is hangsúlyozta.

A Bizottság javaslata kapcsán azonban aggályokat vetett fel, hogy a tilalom hatálya kizárólag a közszektorra korlátozódott. Ahogy arra az Európai Adatvédelmi Testület és az európai adatvédelmi biztos 5/2021. sz. közös véleménye is utalt, a magánvállalatok, például a különböző közösségimédia- és felhőszolgáltatók, a közhivatalokhoz hasonlóan hatalmas mennyiségű személyes adatot kezelnek, mely felhasználásával könnyedén végezhetnek kiterjedt közösségi pontozást.¹¹³ Lényeges szempont továbbá az is, hogy a magánszektorban tevékenykedő MI technológiákat alkalmazó cégek napjainkban már olyan létfontosságú infrastruktúrákat ellenőriznek, mint a szállítmányozás, a telekommunikáció vagy a közlekedés. Az e szektorokból történő pontozáson alapú negatív elbírálás vagy kizárás legalább annyira súlyos társadalmi-gazdasági következményekkel járhat az egyének számára, mint az állam által nyújtott szolgáltatásokból való kizárás.

A szabályozással kapcsolatban megemlítendő az Európai Gazdasági és Szociális Bizottság felvetése is, miszerint az i. és ii. alpontban foglalt tiltó feltételek megfogalmazása pontosításra szorulnak. A Bizottság javaslatában ugyanis nem könnyű egyértelmű határvonalat húzni a meghatározott célból történő, de megengedett pontozás, és az *általános közösségi pontozásnak* tekinthető esetek között, vagyis aközött, amikor az értékeléshez felhasznált információ még releváns és észszerűen kapcsolódik az értékelés céljához, és aközött, amikor a felhasznált információ már nem tekinthető az értékelés céljához észszerűen kapcsolódónak.¹¹⁴

Végezetül érdemes megemlíteni, hogy több jogvédő szervezet is kritikával illette a rendelet-tervezetet amiatt, hogy az kizárólag az *általános társadalmi pontozást*, illetve a személyek *megbízhatóságát* értékelő pontozási eljárásokat rendeli tiltani. Érvelésük szerint a pontozáson alapuló eljárásoknak több olyan – jelenleg nagy kockázatúként kategorizált – alkalmazási módja is létezik (pl.: szociális ellátáshoz való hozzáférés engedélyezése, jövőbeni

¹¹³ Európai Adatvédelmi Testület (2021): Az Európai Adatvédelmi Testület és az európai adatvédelmi biztos 5/2021. sz. közös véleménye a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról szóló európai parlamenti és tanácsi rendeletre irányuló javaslatról 3.o.

¹¹⁴ Európai Gazdasági és Szociális Bizottság (2021): Vélemény - Javaslat európai parlamenti és tanácsi rendeletre a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról INT/940 2021, 5.o.

bűnelkövetés/bűnismétlés valószínűségének előrejelzése), ami szintén jelentős veszélyeket rejt az állampolgárok alapvető szabadságjogaira nézve. Álláspontjuk szerint az állampolgárok hosszú távú érdekeire tekintettel a jövőben ezek tiltása is szükségessé válhat.¹¹⁵

4.5.3 Biometrikus azonosítási rendszerek valós idejű használata

A rendelet-tervezet főszabályként tiltja a „valós idejű” távoli biometrikus azonosító rendszerek használatát a nyilvánosság számára hozzáférhető helyeken bűnüldözési célokból”.¹¹⁶

Egyes kiemelt bűnüldözési célok tekintetében azonban megengedte e rendszerek használatát. Az ilyen helyzetek közé tartozott a bűncselekmények potenciális áldozatainak (pl.: eltűnt gyermekek) felkutatása; a természetes személyek életét vagy fizikai biztonságát fenyegető veszélyek vagy a terrortámadás veszélye; valamint az 2002/584/IB tanácsi kerethatározatban felsorolt és az érintett tagállam joga szerint *legalább három évig terjedő szabadságvesztéssel büntetendő bűncselekmények elkövetőjének vagy gyanúsítottjának felderítése, azonosítása és büntetőeljárás alá vonása*.¹¹⁷

A biometrikus azonosító rendszerek bűnüldöző hatóságok általi alkalmazásának létjogosultsága régóta viták tárgyát képezi. Ennek oka, hogy a nyilvánosság számára hozzáférhető helyeken történő távoli biometrikus azonosítás az egyének magánéletébe való betolakodás magas kockázatát hordozza magában. Ezek az eszközök azáltal, hogy képesek azonosítani, egyedileg megjelölni, valamint nyomon követni a kijelölt célszemélyt, komoly fenyegetést jelenthetnek az egyén alapvető szabadságjogaira, ideértve többek között a magánélethez és az adatvédelemhez való jogot, a gyülekezési szabadságot, valamint az egyenlőséghez és a megkülönböztetésmentességhez való jogot is.

A javaslat-tervezetben vázolt megközelítés az efféle gyakorlatoknak csak egy töredékét tiltja, hiszen a valós idejű távoli biometrikus azonosító rendszerek ezen technológiacsalád elenyésző

¹¹⁵Access Now (2021): Access Now’s submission to the European Commission’s adoption consultation on the AI Act. 15.o. Elérhető: <https://www.accessnow.org/wp-content/uploads/2021/08/Submission-to-the-European-Commissions-Consultation-on-the-Artificial-Intelligence-Act.pdf> (2022. 05. 09.); Human Rights Watch (2021): How the EU’s Flawed Artificial Intelligence Regulation Endangers the Social Safety Net. 25.o. Elérhető: <https://www.hrw.org/news/2021/11/10/how-eus-flawed-artificial-intelligence-regulation-endangers-social-safety-net> (2022. 05. 10)

¹¹⁶ MI rendelet-tervezet 5. cikk (1) d.

¹¹⁷ MI rendelet-tervezet 5. cikk (1) d i,ii,iii.

részét képviselik. A javaslat a *nem valós idejű*¹¹⁸ és a *megközelítőleg valós idejű* biometrikus azonosítást továbbra is megengedi a hatóságok számára, annak ellenére, hogy a gyakorlatok szabadságsértő jellege nem minden esetben függ attól, hogy az azonosításra valós időben kerül-e sor vagy sem. Könnyen belátható, hogy a bünyügyi hatóságok által véghez vitt nem valós idejű távoli biometrikus azonosítás – például egy tüntetéssel vagy kültéri politikai tiltakozással összefüggésben – megközelítőleg azonos mértékű visszatartó hatással bírhat.¹¹⁹ Az állampolgárok jogai tekintetében szintén problémás lehet, hogy a tiltás hatálya jelen esetben kizárólag a bűnüldöző hatóságokra és a bűnüldözési célokra terjed ki, mely korlátozás lehetővé teszi e rendszerek más célokra és bármely más szereplő általi használatát.¹²⁰

A rendelet-tervezet a biometrikus azonosító rendszerek használatához köthető gyakorlatok többségét *nagy kockázatúnak* minősítette. Ennek aránytalanságát hangsúlyozva az Európai Adatvédelmi Testület, az európai adatvédelmi biztos és az Európai Gazdasági és Szociális Bizottság is lényegesen szigorúbb szabályozást tartott indokoltnak. Véleményükben nyomatékosan szorgalmazták, hogy a nyilvánosság számára hozzáférhető helyeken általános jelleggel tiltsák be (1) az MI-rendszerek automatikus biometrikus felismerésre (pl.: arc, járás, hang, ujjlenyomat, DNS alapján) történő felhasználását (kivétel ha meghatározott körülmények között hitelesítési célt szolgál); (2) azon gyakorlatokat melyek biometrikus adatokat felhasználó MI-rendszerek segítségével egyéneket csoportosítsanak etnikai hovatartozás, nem, politikai vagy szexuális irányultság szerint; (3) az olyan eljárásokat melyekben az MI biometrikus adatok elemzésével egy természetes személy jövőbeli viselkedésének előrejelzésére vagy besorolására törekszik.¹²¹

4.6 Az Európai Unió Tanácsának közös álláspontja

¹¹⁸ A valós idejű azonosító rendszerek élő anyagot - például élő videofelvételt - használnak egy adott személy azonosítására. A nem valós idejű rendszerek esetében ezzel szemben a biometrikus adatokat már rögzítették, és az azonosításra csak később kerül sor. Ebben az esetben tehát az azonosításra alkalmas képek vagy videofelvelelek azok konkrét használat megelőzően keletkeztek.

¹¹⁹ Access Now (2021): Access Now's submission to the European Commission's adoption consultation on the AI Act. 16-17.o.

¹²⁰ A biometrikus rendszerek alkalmazása a magánszektor részéről (pl.: Clearview AI) szintén kockázatot jelent az egyénekre és a társadalomra nézve. Ráadásul bizonyos esetekben még annál is korlátozottabb jogorvoslati lehetőség áll a sértett(ek) rendelkezésre, mint amikor ezeket a technológiákat hatósági szervek használják.

¹²¹ Európai Gazdasági és Szociális Bizottság (2021): Vélemény - Javaslat európai parlamenti és tanácsi rendeletre a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról INT/940, 5-6o., Európai Adatvédelmi Testület (2021): Az Európai Adatvédelmi Testület és az európai adatvédelmi biztos 5/2021. sz. közös véleménye. 13-16o.

Az Európai Unió Tanácsa 2022. december 6-án tette közzé az MI rendelet-tervezettel kapcsolatos közös álláspontját (általános megközelítés).¹²² Több fontos módosítási javaslat is szerepelt a szövegben.

A Bizottság által bevezetett - és több helyütt is kritikával illetett – általános, kiterjesztő MI definícióhoz¹²³ képest, a Tanács az MI szűkebb definíciójának használatát javasolta, amely az MI-re jellemző sajátosságok közül a gépi tanulásra, valamint a logikai- és tudásalapú fejlesztésekre fektette a hangsúlyt. Ennek célját abban jelölték meg, hogy egyértelműbben meg lehessen különböztetni az MI-t az egyszerűbb szoftverrendszerektől¹²⁴ így módon rögzítve az azok között fennálló különbségeket. A közös álláspont javaslata szerint az MI; *olyan rendszer, amelyet úgy terveztek, hogy autonóm elemeket is alkalmazva működjön, és amely gépi és/vagy ember által szolgáltatott adatok és bemeneti adatok alapján – gépi tanulást és/vagy logikai és tudásalapú megközelítéseket alkalmazva – kikövetkezteti, hogyan érhető el a célkitűzések egy adott csoportja, és a rendszer által generált kimeneteket hoz létre, például olyan tartalmakat (generatív MI-rendszerek), előrejelzéseket, ajánlásokat vagy döntéseket, amelyek befolyásolják azokat a környezeteket, amelyekkel az MI-rendszer kölcsönhatásba lép.*¹²⁵

Lényeges változtatás, hogy a Tanács közös álláspontjában az általános társadalmi pontozás gyakorlatának tilalmát a közszektorról a magánszereplőkre is kiterjesztette, így a II. Cím 5. cikk (c) bekezdésében már nem szerepelt az MI-rendszerek „*hatóságok általi vagy nevükben történő*” korlátozás. Ennek értelmében a közigazgatási szerveken túl, már a magánvállalatok sem alkalmazhatnak olyan MI rendszereket, melyek állampolgárokat *értékelnek vagy osztályoznak közösségi magatartásuk, illetve ismert vagy előre jelzett személyes vagy*

¹²² Az Európai Unió Tanácsa (2022): Javaslat – Az Európai Parlament és a Tanács rendelete a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról – Általános megközelítés. 14954/22,

¹²³ Indoklása szerint az Európai Bizottság *technológiasemleges* és *időtálló* definíció megalkotására törekedett, mely kellően rugalmas ahhoz, hogy alkalmazkodjon az állandó műszaki fejlődéshez, ugyanakkor elég pontos ahhoz, hogy a jövőben is garantálja a szükséges jogbiztonságot. A 2021 áprilisában közzétett rendelet-tervezet meghatározásában az MI olyan szoftver, *amelyet meghatározott technikák és megközelítések alkalmazásával fejlesztettek, és amely az ember által meghatározott célkitűzések adott csoportja tekintetében olyan kimeneteket, például tartalmat, előrejelzéseket, ajánlásokat vagy döntéseket képes generálni, amelyek befolyásolják azt a környezetet, amellyel kölcsönhatásba lépnek.*

¹²⁴ Míg a hagyományos szoftverek általában előre meghatározott utasításokon alapulnak és konkrét feladatok elvégzésére vannak programozva, az MI rendszerek képesek tanulni, alkalmazkodni és döntéseket hozni az adatok elemzésével és azokból levont következtetésekkel. A Tanács érvelése alapján az MI és a hagyományos szoftverrendszerek közötti különbség megértése azért is lényeges, mert így nyílik lehetősége a szabályozóknak az MI specifikus kihívások és lehetőségek felismerésére.

¹²⁵ Általános megközelítés I. Cím 3. cikk (1)

*személyiségi tulajdonságok alapján, oly módon, hogy az így kapott értékelés aránytalanul kedvezőtlen bánásmódhoz vezetne a kínált szolgáltatás tekintetében.*¹²⁶

A Tanács ugyan nem fogadta el a Régiók Európai Bizottságának azon javaslatát, hogy *a szubliminális technikákat* alkalmazó MI gyakorlatokkal összefüggésben terjessze ki a tiltó feltételt a testi és lelki károsodáson túl a *pénzügyi veszteségek és gazdasági diszkrimináció* eseteire is, azonban a magatartást torzító gyakorlatok második nevesített formája esetében már javasolja a tiltás kiterjesztését az életkor, és a szellemi és testi fogyatékosságon túl azon személyekre is, akiknek a sebezhetősége *a szociális vagy gazdasági helyzetükből* fakad. Ennek megfelelően a II. Cím 5. cikkének (b) bekezdése tiltja az olyan MI rendszerek forgalomba hozatalát, üzembe helyezését vagy használatát, amelyek *a személyek egy meghatározott csoportjának életkor, fogyatékosság, illetve egyedi szociális vagy gazdasági helyzet miatt fennálló valamilyen sebezhetőségét kihasználják.*¹²⁷

A Tanács a magatartást torzító gyakorlatokkal kapcsolatban az Európai Gazdasági és Szociális Bizottság és a Régiók Európai Bizottságának azon felvetését sem tartotta megalapozottnak, miszerint testi és lelki károsodást okozó gyakorlatokon túl, szükséges volna azon magatartástorzító technikák tiltása is, melyek hosszú távon alááshatják a *demokráciát és a jogállamiságot.*

A Tanács szintén nem látta indokoltnak a biometrikus azonosítási rendszereket érintő lényegesen szigorúbb szabályozás megalkotását. E rendszerekkel összefüggésben az általános megközelítés két technikai változtatást eszközölt: (1) pontosította azokat a célkitűzéseket, amelyek esetében a valós idejű és távoli azonosítás bűnüldözési célokból szigorúan szükséges, valamint; (2) pontosította a *távoli biometrikus azonosító rendszer* és a *valós idejű távoli biometrikus azonosító rendszert* fogalmait annak tisztázása érdekében, hogy mely helyzetek tartoznak a kapcsolódó tilalom és nagy kockázatú felhasználás körébe, és melyek nem.

A későbbiekben tárgyalandó GenMI modellekhez kötődően fontos kiemelni, hogy a Tanács az átláthatósági követelményeket rögzítő 52. cikk (3) bekezdését is módosította, melynek értelmében az MI által generált (audio- vagy videótartalom) kimenetet nem kell mesterségesen előállítottak címkézni *„ha a tartalom nyilvánvalóan kreatív, satirikus, művészi vagy fikciós*

¹²⁶ Például, ha egy bank vagy biztosító társaság által alkalmazott értékelési rendszer (amelynek egyik célja hogy a pontszámok alapján döntsön az ügyfél biztosítási díjáról), olyan általános kockázati profilt hoz létre, amely ellehetleníti az értékelt személyt bármely pénzügyi szolgáltatásokhoz való hozzáféréshez.

¹²⁷ Általános megközelítés II. Cím 5. cikk (b)

mű vagy program része”.¹²⁸ Ez a (dolgozat témáját illetően is) lényeges változás, - ahogy arra Mezei utalt egy 2023-as tanulmányában - már annak okán is figyelemreméltó, hogy az MI által létrehozott alkotások a jelenlegi európai szerzői jogi rendszerben nem élvezhetnek ilyen védelmet.¹²⁹

4.6.1 Új kategória: Általános célú MI

A Tanács által javasolt kompromisszumos szöveg legfontosabb változtatásának az tekinthető, hogy bevezetett egy teljesen új MI kategóriát is, *általános célú MI* néven. A Tanács által bevezetett új kategóriát többek között a GenMI modellek rohamos fejlődése és tömeges hozzáférhetővé válása indokolta. Ezek a modellek eredeti tartalom létrehozására képesek, illetve számos alkalmazási lehetőséggel rendelkeznek. Fontos jellemzőjük a gépi tanulás; a betáplált nagymennyiségű adat feldolgozásával megtanulják értelmezni az azokban rejlő mintákat, kapcsolatokat és struktúrákat, majd a tanulási folyamat végeztével új (a betáplált adatokban nem fellelhető) tartalom létrehozására válnak alkalmassá. Mivel az általános célú MI-k számos területen alkalmazhatóak, valamint sokkal gyorsabban válhatnak más MI-rendszerek kiegészítő részévé, mint ahogyan a jogalkotási folyamatok erre reagálni tudnának a Tanács által elfogadott közös álláspontban egy olyan definíció megalkotására törekedtek, amely a lehető legszélesebb körben határozza meg az ilyen technológiákat. Ennek értelmében a közös álláspont az általános célú MI-t olyan rendszerként definiálta *„amelyet – függetlenül attól, hogy miként hozták forgalomba vagy helyezték üzembe, ideértve a nyílt forráskódú szoftvert is – a szolgáltató általánosan alkalmazható funkciók, például kép- vagy beszédfelismerés, hang- és videóelőállítás, mintaészlelés, kérdésmegválaszolás, fordítás és egyéb funkciók ellátására szán; az általános célú MI-rendszerek számos összefüggésben használhatók és számos egyéb MI-rendszerbe integrálhatóak.”*¹³⁰

¹²⁸ 52. cikk (3) bekezdés: *Azon MI-rendszerek felhasználói, amelyek olyan, meglévő személyekre, tárgyakra, helyekre vagy más szervezetekre vagy eseményekre érzékelhetően hasonlító képet, audio- vagy videotartalmat generálnak vagy manipulálnak, amely egy személy számára megtévesztő módon eredetinek vagy valóságosnak tűnhet („deepfake”), közlik, hogy a tartalmat mesterségesen hozták létre vagy manipulálták. Az első albekezdés azonban nem alkalmazandó abban az esetben, ha a felhasználást törvény engedélyezi bűncselekmények felderítése, megelőzése, kivizsgálása és eljárás indításának céljából, vagy ha a tartalom nyilvánvalóan kreatív, szatirikus, művészeti vagy fikciós mű vagy program részét képezi, a harmadik felek jogaira és szabadságaira vonatkozó megfelelő biztosítékok mellett.*

¹²⁹ Mezei Kitti (2023): A mesterséges intelligencia jogi szabályozásának aktuális kérdései az Európai Unióban. In Medias Res. 2023/1. 4, 60.o.

¹³⁰ Általános megközelítés I. Cím 3. cikk (1b)

Amennyiben az ilyen MI-rendszereket nagy-kockázatú MI-ként, vagy azok alkotórészeként használják, meg kell felelniük a rendelet III. címének 2. fejezetében meghatározott nagy kockázatú MI-kel kapcsolatos követelményeknek. A közös álláspont 4b. cikkének (1) bekezdése szerint azonban ezen követelmények nem alkalmazhatóak közvetlenül az általános célú MI-rendszerekre. Ehelyett egy részletes hatásvizsgálaton alapuló végrehajtási jogi aktus határozná meg az alkalmazásukat, „*mégpedig jellemzőik, műszaki megvalósíthatóságuk, az MI-értéklánc sajátosságai, valamint a piaci és technológiai fejlődés fényében*”.¹³¹

4.7 A Parlament kompromisszumos javaslata

2023 május 11-én az Európai Parlament Belső Piaci és Fogyasztóvédelmi Bizottsága (IMCO), és az Állampolgári Jogi, Bel- és Igazságügyi Bizottsága (LIBE) is jóváhagyta a Parlament által benyújtott kompromisszumos szövegjavaslatot, amelyet aztán a Parlament 2023. június 14-én fogadott el hivatalosan.

A tiltott gyakorlatok szabályozására vonatkozóan több módosítási javaslat is említésre méltó. A magatartástorzító gyakorlatokkal összefüggésben a Parlament kiterjesztett a szubliminális technikák alkalmazására irányuló tiltást a döntéshozatali képesség gyengítésére alkalmas célzottan manipulatív vagy megtévesztő technikák használatára is.¹³² A kompromisszumos szöveg szigorított a biometrikus azonosító szoftverek alkalmazásával kapcsolatosan és kiterjesztette a tilalmat az ilyen azonosító rendszerek ex-post (tehát utólagos) alkalmazására is.¹³³ Emellett az elfogadhatatlan kockázatú alkalmazásokat tartalmazó 5. cikk új elemmel is bővült a prediktív rendszet keretében történő lehetséges jövőbeni bűnelkövetés és bűnismétlés kockázatának értékelésére vonatkozó gyakorlatok tiltásával.¹³⁴ Szintén új elem, hogy a szöveg

¹³¹ Általános Megközelítés IA. Cím 4b. cikk (1)

¹³² A javaslat 5. cikk 1.bek a pontja tiltani rendeli az olyan MI-t, amelyek *...célzottan manipulatív vagy megtévesztő technikákat alkalmaznak azzal a céllal vagy olyan hatás érdekében, hogy lényegesen torzítsák egy személy vagy személyek egy csoportjának magatartását azáltal, hogy jelentősen gyengítik az adott személy megalapozott döntéshozatalra való képességét, és ennek következtében a személy olyan döntést hozzon, amelyet egyébként nem hozott volna meg, oly módon, amely e személynek, egy másik személynek vagy személyek egy csoportjának jelentős károsodást okoz vagy okozhat;*

¹³³ 5 cikk 1 bekezdés dd pont

¹³⁴ A kompromisszumos szöveg alapján nem engedélyezett az olyan MI alapú kockázatelemző rendszer alkalmazása, mely annak érdekében végez kockázatelemzést, hogy (1) felmérje *egy természetes személy milyen kockázatot jelenthet egy bűncselekmény vagy szabálysértés elkövetése vagy újbóli elkövetése szempontjából,* illetve, hogy (2) *profilalkotás, személyiségjegyek, tulajdonságok, vagy múltbeli bűnöző magatartás értékelése*

tiltani rendelte a természetes személyek érzelmeiből következtetést levonó MI-rendszerek¹³⁵ forgalomba hozatalát, üzembe helyezését vagy használatát *a bűnüldözés, a határigazgatás, a munkahelyek és az oktatási intézmények területén.*¹³⁶ Emellett tiltotta az olyan MI rendszerek használatát is, melyek *természetes személyeket érzékeny vagy védett tulajdonságok vagy jellemzők szerint, vagy ilyen tulajdonságokból vagy jellemzőkből levont következtetések alapján kategorizálnak* (biometrikus kategorizálási rendszer).¹³⁷ A javaslat szintén elfogadhatatlan kockázatúként kategorizálta az internetről vagy CCTV felvételekről származó képek lekérdezésével kiterjedt arcfelismerő adatbázisokat létrehozó rendszereket is.¹³⁸

A nagy kockázatú MI-k tekintetében lényeges változás, hogy míg a Bizottság és a Tanács korábbi javaslatainak III. mellékletében felsorolt területeken alkalmazott modellek automatikusan magas kockázatúnak lettek minősítve, addig a Parlament javaslatában egy további szint került bevezetésre a nagy kockázatú MI kategória feltételeként. Ennek értelmében a mellékletben felsorolt MI modellek csak abban az esetben minősültek nagy kockázatúnak, *ha az egészségre, biztonságra vagy alapvető jogokra nézve jelentős veszélyt jelentenek.*¹³⁹

Emellett bővítésre került azon alkalmazási területek köre is, melyekben az MI nagy kockázatúnak minősülhet.¹⁴⁰ Figyelembe véve az Európai Gazdasági és Szociális Bizottság

alapján előre jelezze a tényleges vagy potenciális bűncselekmények előfordulását vagy megismétlődését. (5. cikk (1) bek. d) pont)

¹³⁵ Az érzelemfelismerő MI-rendszer az *egyének vagy csoportok biometrikus és biometrikus alapú adatai alapján azonosítja vagy levezeti a természetes személyek érzelmeit, gondolatait, lelkiállapotát vagy szándékait* (3. cikk (1) bek. 34. pont)

¹³⁶ Tiltott a természetes személyek érzelmeiből következtetést levonó MI-rendszerek forgalomba hozatala, üzembe helyezése vagy használata a bűnüldözés, a határigazgatás, a munkahelyek és az oktatási intézmények területén. 5. cikk (1) bek. (dc) Az érzelemfelismerő rendszereket illetően bár üdvözlendő korlátozásokat fogalmazott meg a javaslat, több magánjogi és emberi jogi szervezet, mint például az Európai Digitális Jogok (European Digital Rights) és az Access Now, teljes körű tilalmat szorgalmaztak. Ezt többek között azzal indokolják, hogy az ilyen rendszerek alkalmazása sokszor pontatlan lehet, mivel az érzelmek megítélése rendkívül összetett és kultúrától függő folyamat, ami jelentős mértékben eltérhet a gépek által "látott" és értékelt viselkedésformákhoz képest. Emellett az érzelmi felismerésre épülő technológiák alkalmazása adatvédelmi és diszkriminációs kockázatokat is rejt magában, amik sérthetik az egyének személyiségi jogait és szabadságait. A pontatlanságról bővebben Ryan-Mosley, Tate (2023): AI isn't great at decoding human emotions. So why are regulators targeting the tech? MIT Technology Review. Elérhető: <https://www.technologyreview.com/2023/08/14/1077788/ai-decoding-human-emotions-target-for-regulators/> (2023. 09. 04.)

¹³⁷ 5.cikk 1 bekezdés b) pont.

¹³⁸ 5 cikk – 1 bekezdés – d) b) pont *olyan MI-rendszerek, amelyek az arcképek internetről vagy zártláncú televízió felvételekből való, nem célzott lekérdezésével arcfelismerő adatbázisokat hoznak létre vagy ilyeneket bővítenek*

¹³⁹ (6. cikk 1-2 bekezdés, 32 preambulum.)

¹⁴⁰ Például több helyütt is módosult a III melléklet 1 bekezdés 8 pontja, melyben eredetileg az igazságügyi hatóságok vagy nevükben eljáró fél által használt MI rendszerek szerepeltek *a tények és a jog kutatásában és értelmezésében, valamint a jog konkrét tényállásra történő alkalmazásában.* Az új szövegben az említett célok tekintetében már a *közigazgatási szervek* vagy az ő nevükben eljáró fél is említésre került. Ez a gyakorlatban azt jelentheti, hogy a bűncselekményi tényállások megállapításán túl, már bizonyos szabálysértések körülményeinek MI rendszerekkel való felderítése esetén is – főszabályként - nagy kockázatúként kell kategorizálni az alkalmazott MI-t.

2021-ben közzétett véleményét, melyben több ízben is hangot adtak az MI rendszerekhez köthető *általánosabb/össztársadalmi szinten értelmezhető veszélyforrásoknak (szavazási magatartás megváltozásával járó demokratikus vagy jogállami hatások)*¹⁴¹, a kompromisszumos szöveg egy teljesen új nagy kockázatú MI alkalmazási területet is bevezet. Amennyiben teljesül a már említett plusz feltétel – tehát hogy az alkalmazás az egészségre, biztonságra vagy alapvető jogokra nézve jelentős veszélyt jelent – a javaslat szerint automatikusan nagy kockázatúnak kell tekinteni azon MI-rendszereket, *amelyeket választások vagy népszavazások eredményének vagy természetes személyek választások vagy népszavazások alkalmával tanúsított szavazási magatartásának befolyásolására szánnak*.¹⁴²

A kompromisszumos javaslat új kötelezettségeket is előírt a nagy kockázatú rendszerekhez kötődően. A nagy kockázatú MI-rendszerek *alkalmazói* kritikus szerepet játszanak az alapvető jogok védelmében, hiszen helyzetüknél fogva az alkalmazók képesek a legjobban megérteni, hogy konkrétan hogyan használják majd a nagy kockázatú MI rendszert, és tudják azonosítani a fejlesztési szakaszban előre nem látott potenciális jelentős kockázatokat. Az alapvető jogok védelmének hatékony biztosítása érdekében a nagy kockázatú MI rendszerek alkalmazójának *alapjogi hatásvizsgálatot* (Fundamental Rights Impact Assessment, FRIA) kell végeznie az alkalmazás megkezdése előtt.¹⁴³ A hatásvizsgálatot részletes tervnek kell kísérnie, amely ismerteti az alapvető jogokat érintő, legkésőbb az alkalmazás megkezdésének időpontjától kezdve azonosított kockázatok mérséklését elősegítő intézkedéseket vagy eszközöket.

A rendelet-tervezet nem tartalmazott megfelelő intézkedéseket azoknak az egyéneknek a védelmére, akik valószínűleg a leginkább károsan érintettek az MI rendszerek használata által, valamint azon közérdekű szervezetek számára, akiknek fontos szerepe lehet ezeknek az egyéneknek a képviselésében. Az állampolgárok jogvédelmét erősítendő a Parlamenti javaslat

¹⁴¹ Európai Gazdasági és Szociális Bizottság (2021): Vélemény - Javaslat európai parlamenti és tanácsi rendeletre a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról INT/940, 4-5.o

¹⁴² *Nem tartoznak ide az olyan MI rendszerek, amelyek kimenetének a természetes személyek nincsenek közvetlenül kitéve, mint például a politikai kampányok adminisztratív és logisztikai szempontból való megszervezéséhez.* III melléklet (1) bek. 8aa

¹⁴³ A jogszabályszöveghez újonnan csatol 29a cikk a-j értelmében a nagy kockázatú MI-rendszerek üzembe helyezése előtt az alkalmazók kötelezően értékelik a rendszerek hatását a felhasználási környezetben. Az értékelésnek tartalmaznia kell a következő elemeket: rendszer rendeltetésének egyértelmű meghatározása, rendszer felhasználása tervezett földrajzi és időbeli hatókörének meghatározása, rendszer használata által valószínűleg érintett természetes személyek és csoportok, annak ellenőrzése, hogy a rendszer használata megfelel-e az alapvető jogokra vonatkozó releváns uniós és nemzeti jognak, az alapvető jogokra gyakorolt, észszerűen előre látható hatás, a marginalizált személyeket vagy kiszolgáltatott csoportokat valószínűleg érintő bármilyen konkrét kár veszélye, a rendszer használatának észszerűen előrelátható kedvezőtlen hatása a környezetre, részletes terv arról, hogy a kárt és az alapvető jogokra gyakorolt negatív hatást hogyan fogják enyhíteni, az alkalmazó által létrehozott irányítási rendszer, beleértve az emberi felügyeletet, a panaszkezelést és a jogorvoslatot is

kiegészíti a Bizottság kezdeti rendelettervezetét azzal, hogy bevezeti a *panasztételi jogot*, amely lehetővé teszi az egyéneknek és csoportoknak, hogy panaszt nyújtsanak be a nemzeti felügyeleti hatóságokhoz, ha úgy érzik, hogy az MI rendszerek döntései megsértették jogaikat.¹⁴⁴ A javaslat értelmében a nemzeti felügyeleti hatóságok kötelesek tájékoztatni a panaszosokat a felülvizsgálati folyamat során és értesíteni őket a döntés kimeneteléről, beleértve azt is, hogy áll-e rendelkezésre bírósági jogorvoslat.¹⁴⁵ Emellett a törvényjavaslat lehetővé teszi az állampolgárok számára, hogy *magyarázatot* kapjanak az MI által hozott döntésekről, ezáltal növelve az átláthatóságot és elősegítve az érintettek tájékoztatását.¹⁴⁶

A Parlament javaslata az Unió belső piac szétagoltság elkerülése, az új rendelet hatékony és összehangolt végrehajtásának elősegítése, és nem utolsósorban a nemzeti felügyeleti hatóságok, valamint az uniós intézmények, szervek, hivatalok és ügynökségek aktív támogatása céljából egy Európai Unió Mesterséges Intelligenciával Foglalkozó Hivatal (MI-Hivatal) létrehozását is javasolta.¹⁴⁷

4.7.1 Széles körű alkalmazás, általános célú felhasználás

A Parlament által közzétett kompromisszumos szöveg különös hangsúlyt fektetett az általános célú MI-k szabályozási kérdéseire, szigorú követelményeket előírva részükre többek között a használati utasítások részletes dokumentálása, az adatforrások megfelelőségének biztosítása, és a környezeti hatások figyelembevétele terén.

A Parlament által alkotott definíció szerint az általános célú MI olyan rendszer "*amely széles körben használható olyan alkalmazásokra vagy igazítható olyan alkalmazásokhoz, amelyekre nem kifejezetten és szándékosan tervezték*".¹⁴⁸ A fogalomalkotás során - hasonlóan a Tanácshoz

¹⁴⁴ 68a. cikk a természetes személy vagy természetes személyek minden csoportja jogosult arra, hogy panaszt tegyen egy nemzeti felügyeleti hatóságnál – ha megítélése szerint a rá vonatkozó MI rendszer megsérti e rendeletet

¹⁴⁵ 68b cikk minden természetes vagy jogi személy hatékony bírósági jogorvoslatra jogosult a nemzeti felügyeleti hatóság rá vonatkozó, jogilag kötelező erejű döntésével szemben /

¹⁴⁶ A 68c. cikk (Az egyéni döntéshozatal magyarázatához való jog) kimondja, hogy a *az alkalmazó által nagy kockázatú MI-rendszer kimenete alapján hozott döntés által érintett személy, amely döntés olyan joghatásokat vagy őt hasonlóan jelentősen érintő hatásokat vált ki, amelyek megítélése szerint hátrányosan befolyásolják egészségét, biztonságát, alapvető jogait, társadalmi-gazdasági jólétét vagy az e rendeletben meghatározott kötelezettségekből eredő bármely más jogát, a 13. cikk (1) bekezdése szerinti egyértelmű és érdemi magyarázatot kérhet az alkalmazótól az MI-rendszer döntéshozatali eljárásban betöltött szerepéről, a meghozott döntés fő paramétereiről és a kapcsolódó bemeneti adatokról.*

¹⁴⁷ 76 Preambulumkezdés

¹⁴⁸ 3 cikk – 1 bekezdés – 1 d pont

- a Parlament is fontosnak tartotta kiemelni, hogy ezeknek a rendszereknek a felhasználási területei az idő előrehaladtával jóval szélesebbé és változatosabbá válhatnak, mint azt kezdetben tervezték. Ez a jellegzetesség a szabályozás tekintetében azért is érdemel körültekintést, mert a modellek könnyű adaptálhatósága előre nem tervezett környezetekhez és funkciókhoz jelentősen megnehezíti azok jövőbeli potenciális kockázatainak pontos előrejelzését.¹⁴⁹

Az új GenMI modellek fokozódó népszerűsége okán az általános célú MI kategória fogalmi tisztázásán túl a Parlament egy új MI modelltypus bevezetését is szükségesnek tartotta. A jogszabályalkotási folyamat utolsó szakaszában több nézeteltérést is kiváltó ún. *alapmodelleket* (foundation model) a kompromisszumos javaslat szövege úgy határozta meg, mint *az MI-rendszer olyan modellje amelyet széles körű adatokon tanítottak, általános jellegű kimenetre terveztek, és amely széles körben különféle feladatokhoz igazítható*”.¹⁵⁰

4.7.2 Az alapmodellek szolgáltatóinak kötelezettségei és a ChatGPT szabály

Az alapmodellek szolgáltatói számára az új 28b cikkben kerültek meghatározásra azon követelmények, melyeknek az MI forgalmazását vagy üzembe helyezését megelőzően kell megfelelniük.¹⁵¹ Ennek értelmében a szolgáltató köteles igazolni, hogy megfelelő *tervezéssel, teszteléssel és elemzéssel az egészséget, a biztonságot, az alapvető jogokat, a környezetet, a demokráciát és a jogállamiságot érintő, észszerűen előrelátható kockázatok azonosítása, csökkentése és mérséklése* a fejlesztést megelőzően és annak során megfelelő módszerekkel (például független szakértők bevonásával) megtörtént.¹⁵² Ezen túl a javaslat hosszasan sorolja a további kötelezettségeket *a minőségi adatkészletek biztosításával* (28. cikk 2b) a megfelelő szintű *teljesítmény, kiszámíthatóság, értelmezhetőség, korrigálhatóság, biztonság és kiberbiztonság* elérésével (28. cikk 2c), az erőforrásfelhasználás és hulladékcsökkentéssel (28.

¹⁴⁹ Részben ehhez kötődik, hogy a Parlament javaslatának a 28 cikk (1) bekezdés, ba pontja kiterjeszti a nagy kockázatú rendszerek szolgáltatókra vonatkozó kötelezettségeket az MI értéklánc azon szereplőire is, akik jelentős módosítást hajtanak végre egy olyan forgalomba hozott, vagy üzembe helyezett MI rendszeren, amely eredetileg nem számított nagy kockázatúnak. Ez a gyakorlatban azt jelenti, hogy ha egy eredetileg nem nagy kockázatú általános célú MI-rendszeren olyan módosításokat hajtanak végre, amelyek következtében az immár nagy kockázatúvá minősül – ami az általános célú MI-k változatos felhasználási köréből adódóan gyakran előfordulhat - akkor a lényeges átalakítást végző fél számít majd az új szolgáltatónak, és rá vonatkoznak a nagy kockázatú rendszerekre előírt szigorúbb szabályok és kötelezettségek.

¹⁵⁰ 3 cikk 1 bekezdés 1 c pont

¹⁵¹ 28b cikk (1)

¹⁵² 28b cikk (2) Illetve a *fejlesztés után még fennálló, nem mérsékelhető kockázatok dokumentálására* is sor került.

cikk 2d) az *átfogó műszaki dokumentációt és érthető használati utasítással készítésével* (28. cikk 2e), valamint a *minőségirányítási rendszer* létrehozásával kapcsolatban (28. cikk 2f). Ezeket a követelményeket a modell teljes életciklusán keresztül független szakértők által kell tesztelni, dokumentálni és ellenőrizni.¹⁵³

Az alapmodellekre vonatkozó általános kötelezettségeken túl, a javaslat az átláthatósággal, a tanítóadatokkal és a szerzői joggal összefüggésben három további jogkötelezettségeket állapít meg a GenMI¹⁵⁴ rendszerekben használt vagy GenMI rendszerekre specializált alapmodellek szolgáltatói számára (ChatGPT szabály).¹⁵⁵ A 28b cikk (4) bekezdése kimondja, hogy a szolgáltatók:

*a) eleget tesznek az 52. cikk (1) bekezdésében felvázolt átláthatósági kötelezettségeknek;*¹⁵⁶

b) úgy tanítják be, és adott esetben úgy alakítják ki és fejlesztik ki az alapmodellt, hogy a technika általánosan elismert állásának megfelelően és az alapvető jogok, köztük a véleménynyilvánítás szabadságának sérelme nélkül megfelelő biztosítékokat nyújtsanak az uniós jogot sértő tartalmak generálásával szemben;

*c) a szerzői jogra vonatkozó uniós, nemzeti vagy uniós jogszabályok sérelme nélkül dokumentálják és nyilvánosan hozzáférhetővé teszik a szerzői jog alapján védett tanulóadatok felhasználásának kellően részletes összefoglalóját*¹⁵⁷

Azt követően, hogy 2022 decemberében a Tanács által közzétett közös álláspont majd fél évvel később 2023 júniusában a Parlament által jegyzett kompromisszumos szövegjavaslat is

¹⁵³ 28b cikk (2)c

¹⁵⁴ A 28b cikk (4) egyúttal definícióval is szolgál a GenMI-re. Ez alapján a GenMI olyan MI rendszereket jelöl, amelyek célja, hogy különböző szintű autonómiával generáljanak olyan tartalmakat, mint például összetett szöveg, kép, hang vagy videó („generatív MI”).

¹⁵⁵ Hacker, Philipp, Andreas Engel, Marco Mauer (2023): Regulating ChatGPT and Other Large Generative AI Models. In Proceedings of the Fairness, Accountability, and Transparency (FAccT '23). ACM, New York. 1115.o. <https://doi.org/10.1145/3593013.3594067>.

¹⁵⁶ Az 52 cikk (1) bekezdés már említésre került az előző fejezetben is a Bizottsági rendelet-tervezet ismertetésénél. Az átláthatósági klauzula így szól: *A szolgáltatók biztosítják, hogy a természetes személyekkel való közvetlen interakcióra szánt MI-rendszereket úgy tervezzék meg és fejlesszék ki, hogy az MI-rendszer, maga a szolgáltató vagy a felhasználó kellő időben, egyértelmű és érthető módon tájékoztassa az MI rendszernek kitett természetes személyt arról, hogy MI-rendszerrel áll kapcsolatban, kivéve, ha ez a körülmények és a használat kontextusa alapján nyilvánvaló. Adott esetben e tájékoztatásnak ki kell terjednie arra is, hogy mely funkciók működnek MI-vel, hogy van-e, emberi felügyelet, és hogy ki felelős a döntéshozatali folyamatban, valamint azokra a meglévő jogokra és eljárásokra, amelyek az uniós és a nemzeti jog szerint lehetővé teszik természetes személyek vagy képviselők számára, hogy tiltakozzanak az ilyen rendszerek rájuk vonatkozó alkalmazása ellen, és bírósági jogorvoslatot kérjenek az MI-rendszerek által hozott döntések vagy az általuk okozott károk ellen, beleértve a magyarázatkéréshez való jogukat is.*

¹⁵⁷ 28b cikk (4) bek

elfogadásra került, az EU jogszabályalkotási rendje szerint 2023 nyarán az eljárás a háromszervi egyeztetés szakaszába (ún. trilógus tárgyalásba, vagy háromoldalú tárgyalásba) léphetett. A háromszervi egyeztetés során a Tanács és a Parlament álláspontjait kellett összehangolni, mely folyamatban a Bizottság egyfajta mediátorként léphetett fel az intézmények között.

4.8 Háromoldalú tárgyalások

2023 teléig összesen öt háromoldalú tárgyalásra került sor, az ötödik és egyben utolsó tárgyalás december 6. és 8. között folyt le, melyben a Tanács és az Európai Parlament – egy „maratoni” több mint 12 órás tárgyalást követően - ideiglenes megállapodásra jutott minden szabályalkotási (és politikai) kérdésben, sikeresen lezárva az intézményközi tárgyalásokat. A jogszabályalkotási folyamat utolsó szakaszát végig kísérték az európai jogalkotók és az EU kormányai - kiváltképp Franciaország, Olaszország és Németország - között meghúzódó véleménykülönbségek és feszültségek. Ezek elsősorban abból a félelem fakadtak, hogy a túlzottan szigorú szabályozás visszavetné az Unió esélyeit a technológiai innovációkra és fejlesztésekre kiélezett globális MI versenyben.¹⁵⁸ Bár javaslat több lényeges eleme körül is alakult ki vita (arcfelismerő rendszerek tiltása, MI-hivatal feladat-és jogkörei) különösen megosztónak a javaslat szövegezése az alapmodellekkel kapcsolatban bizonyult. A nagy teljesítményű alapmodellekre vonatkozó követelmények tekintetében az uniós döntéshozók sokáig nem jutottak közös nevezőre. Az őszi háromoldalú tárgyalások alatt egyetértés mutatkozott abban, hogy az alapmodellekre vonatkozó szabályokat többszintű megközelítéssel (layered approach) vezetik be, azaz szigorúbb szabályokat kell alkalmazni a legerősebb, a társadalomra nagyobb hatást gyakorló modellekre.

Azonban a Miniszterek Tanácsának ülésén több tagállam képviselői - az előbb említett Franciaország, Olaszország és Németország vezetésével - szembe helyezkedtek az alapmodellekre vonatkozó szabályozási elképzeléssel.¹⁵⁹ Erre válaszul az Európai Parlament több tisztviselője is kísértelt egy későbbi ülésről, világossá téve, hogy az alapmodellek szigorúbb

¹⁵⁸ Bertuzzi, Luca (2023): EU's AI Act negotiations hit the brakes over foundation models. EURAVTIV. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/eus-ai-act-negotiations-hit-the-brakes-over-foundation-models/> (2024. 01. 24.)

¹⁵⁹ Európa három legnagyobb gazdasága azért is lobbizott, hogy az alapmodellek terén kérte, a kötelezettségek fókuszát önkéntesen vállalt magatartási kódexekre szűkítsék, (többek között mert el akarták kerülni az európai technológiai innováció és kiváltképp az olyan ígéretes európai startupokat, mint a Mistral AI és az Aleph Alpha korlátozását).

szabályozásának mellőzése politikailag nem elfogadható alternatíva.¹⁶⁰ Ha decemberben nem született volna megállapodás az alapmodellekkel kapcsolatos megközelítés általános újragondolásának igénye szinte lehetetlenné tette volna, hogy az EU választások előtt elfogadhassák a jogszabályt mely könnyen ahhoz vezethetett volna, hogy az EU elveszíti előnyét más szabályozó törekvésekkel és szabályozni kívánó országokkal szemben. Végül az akkor soros spanyol elnökség végül december 11-én bejelentette a kompromisszum elérését.

A legnagyobb változás, hogy Tanács általános megközelítésében és a Parlament kompromisszumos javaslatában vázolt egységes szabályozás helyett a decemberi megállapodás két különböző típusú általános célú MI-rendszert határozott meg, amelyekre eltérő szabályok és jogi kötelezettségek vonatkoztak. Az első kategóriát a normál alapmodellek alkotják, melyek alacsonyabb és kevésbé kiterjedt kockázatot hordoznak, míg a második kategóriát az *un. rendszerszintű kockázatot jelentő alapmodellek* alkották.

Az új szabályozási kétszintű szabályozási keretrendszerben az első szint minden általános célú MI-re vonatkozik.¹⁶¹ Ezen a szinten a szabályok olyan követelményeket írnak elő, mint a betanítási módszerekkel kapcsolatos információk nyilvánosságra hozatala, az energiafelhasználás transzparens bemutatása, valamint a szerzői jogi törvények betartásának biztosítása.¹⁶² A második szint kizárólag az *un. nagy hatású rendszerszintű kockázatot jelentő alapmodellekre* vonatkozott. Az ilyen modellek szolgáltatóira szigorúbb kötelezettségek vonatkoztak, amelyek magukban foglalják a modellek értékelését, a rendszerszintű kockázatok felmérését és mérséklését, az ilyen kockázatokról szóló jelentéstételt a Bizottságnak, támadhatósági tesztek végrehajtását, valamint a magas szintű kiberbiztonsági védelem biztosítását.

Mivel a decemberi megállapodás, a belga soros elnökség által közölt 2024 január 26-ai módosítások és az Európai Parlament által március 13-án elfogadott végleges szöveg között néhány hónap telt csak el és a jogszabályi szövegek között inkább nyelvtani és csak kisebb mértéken előforduló technikai különbségek vannak, az érthetőség és egyszerűbb követhetőség

¹⁶⁰ Bertuzzi, Luca (2023): France, Germany, Italy push for ‘mandatory self-regulation’ for foundation models in EU’s AI law. EURACTIV. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/france-germany-italy-push-for-mandatory-self-regulation-for-foundation-models-in-eus-ai-law> (2024. 01. 20.)

¹⁶¹ Ezek alól csak azok a rendszerek képeznek kivételt, amelyek kizárólag kutatási célokra szolgálnak, vagy nyílt forráskódú licenc alatt kerülnek kiadásra?

¹⁶² Az alapmodell fejlesztőinek kötelező nyilvánosságra hozni azon szerzői jog által védett (főként szöveges) adattartalmakat, amelyeket az MI modellek tanításához használtak. A tanítóadatok és a szerzői jog kérdéskörei kiténtetett figyelmet érdemel a GenMI rendszerek kihívásai területén ezért a dolgozat későbbi részében még bővebb kifejtésre kerül.

kedvéért kizárólag végső szöveg rövid bemutatásával és néhány következtetés levonására összpontosít a fejezet befejező része.

4.9 A 2023 március 13-án elfogadott törvényszöveg

XXX- XXX

4.9.1 Mi az MI?

Az MI végleges definíciója kisebb változtatásokon átesett a Parlamenti 2023 nyári kompromisszumos javaslata óta, de lényegében megmaradt az OECD által is ajánlott értelmezési paradigmában, és céljaként azt tűzte ki, hogy megkülönböztesse az MI rendszereket az olyan hagyományos szoftverrendszerektől, amelyek tipikusan előre meghatározott szabályok alapján működnek, és nem képesek önálló tanulásra vagy adaptációra. Az Unió jogalkotó szervek emellett az is hangsúlyozták, hogy a megalkotott fogalmat szorosan össze kell hangolni a MI-vel foglalkozó nemzetközi szervezetek munkájával a *jogbiztonság szavatolása*, valamint a *nemzetközi konvergencia* és a *széles körű elfogadottság előmozdítása* érdekében.¹⁶³ E szempontokat figyelembe véve a jogszabály végleges szövege úgy határozza meg az MI-t, mint olyan: *gépi alapú rendszer, amelyet különböző autonómiaszinteken történő működésre terveztek, és amely a bevezetését követően alkalmazkodóképességet tanúsíthat, és amely a kapott bemenetből – explicit vagy implicit célok érdekében – kikövetkezteti, miként generáljon olyan kimeneteket, mint például előrejelzéseket, tartalmakat, ajánlásokat vagy döntéseket, amelyek befolyásolhatják a fizikai vagy a virtuális környezetet.*¹⁶⁴

4.9.2 Kockázatos gyakorlatok

A továbbiakban három kategória rövid ismertetésére került sor: (1) az elfogadhatatlan kockázatot jelentő– ennek okán tiltott gyakorlatok (2) a nagy kockázatú MI-k (melyekkel

¹⁶³ MI-rendelet Preambulum (12)

¹⁶⁴ MI-rendelet 3. cikk (1)

kapcsolatos szabályok a törvényszöveg körülbelül kettőharmadát fedik le) és az utólag új kategóriaként megalkotott (3) általános célú MI.

Tiltott MI-gyakorlatok

Az Európai Bizottság által 2021. áprilisában közzétett rendelet-tervezet négy olyan elfogadhatatlan kockázattal bíró MI gyakorlatot nevesített, melynek tiltása elengedhetetlen az uniós polgárok alapvető jogainak hosszútávú védelme és biztosítása érdekében. Az előbbieken ismertetett szabályalkotási eljárás során e kezdeti listát végül további gyakorlatokkal is kiegészítették, így a végleges szöveg nyolc különböző alkalmazáskategóriát rendel tiltani. Ezek a gyakorlatok az egyének döntéshozási autonómiájának korlátozására, az általános társadalmi pontozásra, a prediktív rendészetre, arcfelismerő adatbázisok létrehozására, a munkahelyen és oktatási intézményben történő érzelmfelismerésre, biometrikus kategorizálásra és a valós idejű távoli biometrikus azonosításra irányulnak.

Az MI-rendelet 5. fejezetében foglaltak értelmében, tilos az olyan MI-rendszerek forgalmazása, üzembe helyezése, vagy használata amelyek:

(1) *szubliminális technikákat vagy célzottan manipulatív illetve megtévesztő technikákat alkalmaznak azzal a céllal, hogy lényegesen torzítsák egy személy vagy személyek egy csoportjának magatartását azáltal, hogy jelentősen gyengítik a megalapozott döntéshozatalra való képességüket...oly módon, amely az említett személynek, egy másik személynek vagy személyek egy csoportjának jelentős károsodást okoz vagy ésszerű valószínűséggel okozhat;*¹⁶⁵

(2) *életkor, fogyatékosság, illetve egyedi szociális vagy gazdasági helyzet miatt fennálló valamilyen sebezhetőséget használják ki azzal a céllal vagy hatással, hogy lényegesen torzítsák egy személy vagy egy csoporthoz tartozó valamely személy magatartását oly módon, amely az érintetteknek jelentős kárt okoz vagy ésszerű valószínűséggel okozhat*

(3) *természetes személyeket vagy csoportokat értékelnek vagy osztályoznak közösségi magatartásuk, illetve ismert, kikövetkeztetett vagy előre jelzett személyes tulajdonságaik vagy személyiségjegyeik alapján, kedvezőtlen bánásmódot eredményezve, vagy olyan kontextusokban, amelyek nincsenek összefüggésben az eredeti adatgyűjtés kontextusával.*¹⁶⁶

¹⁶⁵ MI-rendelet

¹⁶⁶ A tilalom nem érinti a természetes személyek azon jogszerű értékelési gyakorlatát, amelyet konkrét célból, az uniós és a nemzeti joggal összhangban végeznek Preambulum 31.

(4) természetes személyek kockázatértékelését végzi annak érdekében, hogy – kizárólag a természetes személyekre vonatkozó profilalkotás vagy személyiségjegyeik és tulajdonságaik értékelése alapján¹⁶⁷ – felmérje vagy előre jelezze annak kockázatát, hogy egy adott természetes személy bűncselekményt követ el

(5) az arcképek internetről vagy zártláncú televízió-felvételekből való, nem célzott lekérdezésével arcfelismerő adatbázisokat hoznak létre vagy ilyeneket bővítenek;

(6) természetes személyek érzelmeiből vonnak le következtetéseket munkahelyeken vagy oktatási intézményekben¹⁶⁸

(7) természetes személyeket biometrikus adataik alapján egyénileg kategorizálnak, hogy ezáltal levezessék vagy kikövetkeztessék faji hovatartozásukat, politikai véleményüket, szakszervezeti tagságukat, vallási vagy világnézeti meggyőződésüket, szexuális életüket vagy szexuális irányultságukat¹⁶⁹

(8) továbbra is fennáll a tiltás azon valós idejű távoli biometrikus azonosító rendszerek használatával kapcsolatosan melyeket a nyilvánosság számára hozzáférhető helyeken bűnüldözési célokból alkalmaznak¹⁷⁰

Nagy-kockázatú MI

A Bizottság által eredetileg meghatározott nyolc kiemelt kockázatot jelentő fő terület megmaradt a rendelet végső változatának III. mellékletében is, igaz lényeges módosításokat

¹⁶⁷ A prediktív rendszert ezen formája az Unió alapvető értékei alapján elfogadhatatlan gyakorlatnak minősül, hiszen az *ártatlanság vélelmével* összhangban természetes személyekre vonatkozó döntés minden esetben kizárólag e személyek tényleges magatartása alapján hozható (és soha nem alapulhat kizárólag a rájuk vonatkozó profilalkotás, személyiségjegyek vagy tulajdonságok – mint például *állampolgárság, születési hely, tartózkodási hely, gyermekek száma, adósságsszint vagy gépjármű típusa* – értékelésén MI-rendelet preambulum (42)

¹⁶⁸ A végleges szövegbe nem került be a Parlament kompromisszumos javaslatába foglalt azon kiterjesztés, amely az érzelmefelismerő MI-k alkalmazását a munkahelyeken és az oktatási intézményeken túl a *bűnüldözés* és a *határigazgatás* területein is tiltani rendelte volna. Az MI-rendelet preambulumának (19) bekezdése alapján az érzelmefelismerő rendszer olyan MI, amely arra szolgál, hogy biometrikus adataik alapján azonosítsa vagy kikövetkeztesse természetes személyek érzelmeit vagy szándékait. A fogalom érzelmekre vagy szándékokra utal, például *boldogság, szomorúság, düh, meglepettség, undor, zavar, izgatottság, szégyen, megvetés, elégedettség és derű*.

¹⁶⁹ E tilalom nem terjed ki a más célok elérése érdekében (pl.: bűnüldözés) jogszerűen beszerzett adatkészletek biometrikus adatok szerint történő címkézésére, szűrésére vagy kategorizálására (pl.: képek hajszín vagy szemszín szerinti válogatására).

¹⁷⁰ Az MI-rendelet preambulumának (32) bekezdése két lényeges indokot említ a valós idejű biometrikus azonosítás szigorúbb szabályozása mellett: (1) az ilyen MI rendszerek technikai pontatlansága torzított eredményekhez vezethet, és diszkriminatív hatásokat eredményezhet - különösen az életkor, az etnikai és a faji hovatartozás, a nem vagy a fogyatékosságok tekintetében (2) a valós időben működő rendszerek használatával kapcsolatos ellenőrzések vagy korrekciók korlátozott módon és mértékben állnak csak rendelkezésre (az ilyen azonosítás "azonnali jelle" okán)

eszközöltek többek között az 1. (biometria), 6. (bűnüldözés) 7. (migráció, menekültügy és határigazgatás) és 8. (igazságszolgáltatás és demokratikus folyamatok) pontban. Ehelyütt csak példálózó jelleggel ismertetve, az MI-rendelet 6. cikk (2) bekezdése szerinti a következő rendszerek és gyakorlatok minősülnek nagy kockázatúnak

1. biometria (például a nem valós idejű távoli biometrikus azonosító rendszerek¹⁷¹; biometrikus kategorizálásra használt rendszerek (nem faji hovatartozás, politikai vélemény, szakszervezeti tagság, vallási vagy világnézeti meggyőződés, szexuális élet vagy szexuális irányultság levezetésére használt); biometrikus érzelemfelismerő rendszerek (nem munkahelyen vagy oktatási intézményben alkalmazott))

2. kritikus infrastruktúra és a kritikus digitális infrastruktúrák (pl.:a közúti forgalom, víz-, gáz-, fűtési vagy villamosenergia-szolgáltatás irányításában és működtetésében használt MI-rendszerek.

3. oktatás és szakképzés (pl.: oktatási intézményekbe való bejutás vagy felvétel meghatározásához használt rendszerek, tanulási eredmények értékeléséhez való használatra szánt MI-rendszerek)

4. foglalkoztatás, a munkavállalók irányítása és az önfoglalkoztatáshoz való hozzáférés (jelentkezők felvételéhez vagy kiválasztásához használt MI, pl.: jelentkezők elemzése és szűrése, jelöltek értékelése, célzott álláshirdetések elhelyezése)

5. alapvető magán- és közszolgáltatásokhoz és ellátásokhoz való hozzáférés és ezek igénybevétele (pl.: természetes személyek állami segítségnyújtási ellátásokra és szolgáltatásokra (lásd: egészségügyi szolgáltatás) való jogosultságának értékelése; természetes személyek hitelképességének értékeléséhez vagy hitelpontszámuk megállapításához való használatra szánt MI-rendszerek¹⁷²

6 bűnüldözés (pl.: poligráfként használt MI; annak értékelésére szolgáló MI, hogy egy természetes személy esetében fennáll-e annak kockázata, hogy bűncselekmény áldozatává válik; annak értékelésére szolgáló rendszer, hogy egy természetes személy esetében fennáll-e

¹⁷¹ Kivételt képeznek azok a rendszerek, amelyek kizárólagos célja annak megerősítése, hogy egy adott természetes személy azonos azzal a személlyel, akinek állítja magát. III. melléklet (1); A rendelet III. fejezet 26. cikk (10) értelmében, a nem valós idejű távoli biometrikus azonosításra szolgáló, nagy kockázatú MI-rendszert alkalmazó személynek előzetesen vagy indokolatlan késedelem nélkül, de legkésőbb 48 órán belül *engedélyt kell kérnie az adott rendszer használatára valamely igazságügyi hatóságtól vagy közigazgatási hatóságtól*. Az ilyen rendszereket kizárólag célzott bűnüldözési célokra lehet használni (pl.: eltűnt személy felkutatása).

¹⁷² Kivételt képeznek a dolgozat második fejezetében is említett -és magyar NAV által is alkalmazott (RADAR) - olyan MI-rendszerek, melyeket pénzügyi csalás észlelésére használnak. III. Melléklet 5b

bűncselekmény elkövetésének vagy ismételt elkövetésének kockázata (ez jelen esetben olyan prediktív rendészeti gyakorlatot jelent, amely nem kizárólag a természetes személyekre vonatkozó profilalkotás¹⁷³ vagy személyiségjegyeik és tulajdonságaik értékelése alapján történik)

7. migráció, menekültügy és határigazgatás (pl.: természetes személyek felderítése, felismerése vagy azonosítása céljából használt MI; a valamely tagállam területére belépni szándékozó vagy oda belépett természetes személy által jelentett kockázat – *többek között biztonsági kockázat, irreguláris migrációs kockázat vagy egészségügyi kockázat* – értékelését végző MI rendszer)

8. igazságszolgáltatás és demokratikus folyamatok¹⁷⁴ (pl.: olyan MI-rendszerek, amelyek célja, hogy segítsék az igazságügyi hatóságokat a ténybeli és a jogi elemek kutatásában és értelmezésében; választások vagy népszavazások eredményének vagy a természetes személyek választásokon vagy népszavazásokon tanúsított szavazási magatartásának befolyásolásához való használatra szánt MI-rendszerek)

Ugyan a törvényben megállapított kötelezettségek hosszabb kifejtésére ehelyütt nincs mód, két lényeges elem kiemelése mégis indokolt; (1) a jogszabály különböző mértékű kötelezettségeket ró az MI-értékláncban részt vevő szereplőkre figyelembe véve azok eltérő felelősségét és befolyását az MI-rendszerek biztonságos fejlesztése, működtetése és alkalmazása terén (2) a Bizottság által közzétett rendelettervezet eredeti elképzelésével szemben, nem minden esetben válik automatikusan nagy kockázatúvá egy MI-rendszer ha az ismertetett 8 pont valamelyik explicit módon felsorolt területén alkalmazták (3) az alapjogi hatásvizsgálat követelményeivel kapcsolatos enyhítések

A korábbi változatokhoz hasonlóan a nagy kockázatú MI-hez kötődően az általános kötelezettségek közé tartozik a kockázatkezelési rendszer létrehozása (9. cikk), magas minőségű adatkormányzás (10. cikk), műszaki dokumentáció (11. cikk), folyamatos naplózás és nyilvántartás (12. cikk), átláthatóság és az alkalmazóknak nyújtott megfelelő tájékoztatás

¹⁷³ Ahogy az előző alfejezetben említésre került, az MI rendelet alapján a kizárólagosan profilalkotáson alapuló prediktív rendszert tilos. (GDPR 3. cikkének 4. pontja ad pontos meghatározást a profilalkotásról, amely a *személyes adatok automatizált kezelésének bármely olyan formája, amelynek során a személyes adatokat valamely természetes személyhez fűződő bizonyos személyes jellemzők értékelésére, különösen a munkahelyi teljesítményhez, gazdasági helyzetéhez, egészségi állapothoz, személyes preferenciákhoz, érdeklődéshez, megbízhatósághoz, viselkedéshez, tartózkodási helyhez vagy mozgáshoz kapcsolódó jellemzők elemzésére vagy előrejelzésére használják*

¹⁷⁴ Megemlítendő, hogy a végleges szövegbe a Parlament azon javaslata sem került elfogadásra, amely az igazságszolgáltatás és demokratikus folyamatok ponton belül nagy kockázatúként kategorizálta volna az olyan online óriásplatformnak által *saját ajánlórendszereikben* használt MI-rendszereket, melyek célja, hogy a szolgáltatás igénybe vevőjének a platformon elérhető, felhasználók által létrehozott tartalmakat ajánljanak

(13. cikk), kötelező emberi felügyelet (14. cikk), pontosság, stabilitás és kiberbiztonság követelményeinek betartása (15.cikk).¹⁷⁵

Ezen felül a III. Fejezet kiemeli, hogy a nagy kockázatú MI rendszerek szolgáltatóinak¹⁷⁶ minőségirányítási rendszert kell bevezetniük (17. cikk), részletes dokumentációt kell készíteniük, melyet a rendszer üzembehelyezését követően 10 évig az illetékes nemzeti hatóságok számára elérhetővé kell tenniük (18. cikk), legalább hat hónapig meg kell őrizniük a rendszereik által automatikusan generált naplót legalább (19. cikk), amennyiben rendszerük nem felel meg a jogszabályban foglalt kötéleteseteknek azonnali korrekciós intézkedéseket kell hozniuk, melyről tájékoztatási kötelezettség is terheli őket (20.cikk).

Az alkalmazók¹⁷⁷ kötelezettségei közé tartozik, hogy biztosítaniuk kell, hogy a nagy kockázatú MI-rendszert a gyártó utasításai szerint használják, és az rendszer rendeltetésszerűen működjön; emberi felügyeletet biztosítsanak az MI rendszer működése során; haladéktalanul *tájékoztatniuk* kell a szolgáltatót vagy az importőrt és az illetékes hatóságokat, ha bármilyen kockázatot vagy nem megfelelőséget észlelnek a rendszerrel kapcsolatban; együtt kell működniük a szolgáltatóval és a hatóságokkal a rendszerrel kapcsolatos bármilyen probléma megoldásában és a szükséges korrekciós intézkedések végrehajtásában; meg kell őrizniük az adott nagy kockázatú MI-rendszer által automatikusan generált naplókat - amennyiben az ilyen naplók az ellenőrzésük alatt állnak - *legalább hat hónapos* időtartamig.¹⁷⁸ A forgalmazóknak¹⁷⁹ ellenőrizniük kell, hogy a nagy kockázatú MI-rendszeren fel van-e tüntetve *az előírt CE jelölés*; amennyiben a birtokukban lévő információk alapján úgy találják, hogy a rendszer nem felel meg a követelményeknek, nem forgalmazhatják azt, amíg a megfelelőséget nem biztosítják; az illetékes nemzeti hatóság indokolt kérésére az ilyen rendszerek forgalmazóinak a hatóság rendelkezésére kell bocsátaniuk a tevékenységeikre vonatkozó valamennyi olyan információt és dokumentációt, amely szükséges annak igazolásához, hogy az említett rendszer megfelel.¹⁸⁰

¹⁷⁵ MI rendelet III. Fejezet 9-15. cikk

¹⁷⁶ A szolgáltató olyan természetes vagy jogi személy, hatóság, ügynökség vagy egyéb szerv, aki vagy amely MI-rendszert vagy általános célú MI-modellt fejleszt vagy fejlesztet, és a saját neve vagy védjegye alatt – akár fizetés ellenében, akár ingyenesen – az MI-rendszert vagy az általános célú MI-modellt forgalomba hozza, vagy az MI-rendszert üzembe helyezi MI-rendelet 3. cikk (3)

¹⁷⁷ Az alkalmazó olyan természetes vagy jogi személy, hatóság, ügynökség vagy egyéb szerv, aki vagy amely a felügyelete alá tartozó MI-rendszert használja, (kivéve, ha az MI rendszert személyes, nem szakmai jellegű tevékenység során használják) MI-rendelet 3.cikk (4)

¹⁷⁸ MI-rendelet III. fejezet 26. cikk

¹⁷⁹ A forgalmazó az a szolgáltatótól vagy az importőrtől eltérő természetes vagy jogi személy az ellátási láncban, aki vagy amely az uniós piacon MI-rendszert forgalmaz. MI-rendelet 3.cikk (7)

¹⁸⁰ MI-rendelet III. fejezet 24. cikk

A Bizottság rendelet-tervezete és a Tanács általános álláspontja még automatikus nagy kockázatúként sorolta az MI rendszereket, ha azokat a III. mellékletben felsorolt területek egyikén szándékoznak használni. A 2023 júniusi Parlamenti kompromisszumos szöveg módosítására építve az elfogadott törvényszöveg 6. cikk (3) bekezdése kimondja, az MI rendszereket annak ellenére sem tekintik automatikusan magas kockázatúnak, „hogy a felsorolt magas kockázatú területek egyikében használják fel, *amennyiben nem jelent jelentős károkozó kockázatot természetes személyek egészségére, biztonságára vagy alapvető jogaira nézve, többek között azáltal, hogy nem befolyásolja lényegesen a döntéshozatal kimenetelét.* A végső változatban szereplő kiegészítés értelmében ez a kivétel akkor alkalmazható, ha a rendszer rendeltetése „*jól körülhatárolt eljárási feladat ellátása*”¹⁸¹, „*egy korábban elvégzett emberi tevékenység eredményének a javítása*”¹⁸², döntéshozatali minták észlelése, *anélkül, hogy helyettesítené vagy befolyásolná a korábban elvégzett emberi értékelést*”¹⁸³ vagy *előkészítő feladat végrehajtása*”¹⁸⁴.¹⁸⁵

Az ilyen rendszerek szolgáltatóinak csupán dokumentálniuk kell „csökkentett kockázatú” értékelésüket, ezt a dokumentációt kérésre továbbítaniuk a hatóságnak (6. cikk (2b) bekezdés) és regisztrálniuk kell a rendszerüket egy nyilvánosan elérhető adatbázisban (51. cikk, 60. cikk).

Az Európai Parlament által 2023-ban kezdeményezett kötelező alapjogi hatásvizsgálatokkal (FIRA) kapcsolatosan megemlítendő, hogy a végső szöveg értelmében csak azon nagy kockázatú MI-rendszerek alkalmazóinak kell elvégezniük azokat, akik *közjogi szervek vagy közszolgáltatásokat nyújtó magánszervezetek*, illetve olyan MI-rendszert használnak, amely *hitelminősítésre vagy biztosítási kockázatértékelésre szolgál*.¹⁸⁶ Ennek következtében a magánszektor túlnyomó többsége mentesülhet e kötelezettség alól.¹⁸⁷

¹⁸¹ 6. cikk (3)a

¹⁸² 6. cikk (3)b)

¹⁸³ 6. cikk (3)c

¹⁸⁴ 6. cikk (3)d

¹⁸⁵ Megemlítendő azonban, hogy ha egy MI rendszer természetes személyek profilozására szolgál, mindig magas kockázatúnak minősül, függetlenül a fent említett kivételektől

¹⁸⁶ A 27. cikk (1) bekezdése kimondja, hogy *azon alkalmazóknak, amelyek közjogi szervek vagy közszolgáltatásokat nyújtó magánszervezetek, valamint a III. melléklet 5. pontjának b) és c) alpontjában említett nagy kockázatú MI-rendszerek alkalmazóinak értékelniük kell azon hatást, amelyet az ilyen rendszerek használata az alapvető jogokra gyakorolhat.* A (2) bekezdés tisztázza, e kötelezettség a nagy kockázatú MI-rendszer *első használatára* vonatkozik. Ha az alkalmazó olyan rendszert használ, amely már átesett a vizsgálaton – és elemei nem estek át érdemi változáson – *támaszkodhat a korábban elvégzett alapjogi hatásvizsgálatokra vagy a szolgáltatók által már elvégzett, meglévő hatásvizsgálatokra.*

¹⁸⁷ Az októberi háromoldalú tárgyalások során a Tanács szorgalmazta ezt az enyhítést. Bővebben: Bertuzzi, Luca (2023): AI Act: EU countries mull options on fundamental rights, sustainability, workplace use. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/ai-act-eu-countries-mull-options-on-fundamental-rights-sustainability-workplace-use/> (2024. 02. 25.)

4.10 Két szintű szabályozási megközelítés.

Első szint (horizontális kötelezettségek) minden általános modell szolgáltatójának meg kell felelnie néhány alapvető követelménynek. Ezeket az 53. cikk (1) bekezdés sorolja: a) el kell készíteniük és naprakészen kell tartaniuk a modell műszaki dokumentációját, beleértve annak tanítási és tesztelési folyamatát, valamint értékelésének eredményeit abból a célból, hogy az MI-hivatal és az illetékes nemzeti hatóságok rendelkezésére bocsássák; b) információkat és dokumentációt kell kidolgozniuk, naprakészen tartaniuk és rendelkezésre bocsátaniuk az MI-rendszerek azon szolgáltatói részére, amelyek az általános célú MI-modellt be kívánják építeni MI-rendszereikbe; c) a szerzői és kapcsolódó jogokra vonatkozó uniós jognak való megfelelésre irányuló politikát kell bevezetniük; d) kellően részletes összefoglalót kell készíteniük – az MI-hivatal által rendelkezésre bocsátott sablonnak megfelelően – és közzétenniük az általános célú MI-modell tanításához használt tartalomról¹⁸⁹

Az 53 cikk (1) bekezdés b pontja tekintetében helyesen ismerték fel az Unió jogalkotók, hogy az általános célú MI-modellek szolgáltatói kiemelt felelősséget viselnek az MI-értéklánc mentén, azáltal, hogy az általuk biztosított modellek számos rendszer alapját képezhetik majd (ezáltal ők downstream szolgáltatóvá válnak).¹⁹⁰ Ha egy MI-rendszer szolgáltatója integrálni akar egy általános MI-modellt (például egy rubik kockákat áruló weboldalon lévő chatbot szolgáltatója integrálja rendszerébe a GPT nagy nyelvi modellt) ő, mint szolgáltató szükséges, hogy tisztában legyen az integrálandó általános modell működésével. Ezért az általános modell

¹⁸⁸ Az általános célú MI széles körben alkalmazhatók különböző feladatokra, akár önállóan, akár más MI-rendszerekbe integrálva. Ha egy általános célú MI modellt egy másik MI rendszerbe integrálnak, (vagy ha ez a modell egy másik rendszer részét képezi), az egész MI rendszert általános célú MI-nek kell tekinteni, amennyiben az integráció révén a rendszer képes változatos célok szolgálatára. (MI-rendelet 100. Preambulum)

¹⁸⁹ Az általános célú MI-modellek előtanítása és tanítása során felhasznált adatokat – köztük a szerzői jog által védett szövegeket és adatokat – illető átláthatóság növelése érdekében helyénvaló, hogy az ilyen modellek szolgáltatói kellően részletes összefoglalót készítsenek és tegyenek nyilvánosan hozzáférhetővé az általános célú MI-modell tanításához használt tartalomról.

¹⁹⁰ 3. cikk (68): a downstream szolgáltató: olyan MI-rendszer – ideértve az általános célú MI-rendszert is szolgáltatója, amelybe MI-modellt integráltak, függetlenül attól, hogy a szolgáltató által biztosított, vertikálisan integrált MI-modellről vagy egy másik szervezet által szerződéses viszonyok alapján biztosított MI-modellről van-e szó.

szolgáltatójának – az 53 cikk (1) a pontjában foglalt átláthatósági intézkedéseken túl – szükségszerű részletes információval szolgálni az integrálni tervező downstream szolgáltatóknak is az MI-modellről.

Rendszerszintű kockázat¹⁹¹

51. cikk alapján az általános célú MI-t akkor kell rendszerszintű kockázatúként besorolni, ha a modell, nagy hatású képességekkel rendelkezik (megfelelő technikai eszközök és módszertanok – többek között mutatók és referenciaértékek – alapján). Automatikusan nagy hatásúnak tekinthető, ha a tanítási folyamathoz szükséges számítási teljesítmény meghaladja a 10^{25} lebegőpontos műveletet. A Bizottság fel van hatalmazva arra, hogy módosítsa a küszöbértékeket és referenciaértékeket, reagálva a technológiai fejlődésre.

Rendszerszintű kockázatot jelentő MI szolgáltatóinak többletkötelezettségei

Az 55. cikk tartalmazza a rendszerszintű kockázatot jelentő általános célú MI szolgáltatóinak többletkötelezettségeit, melyek (1) Modell rendszeres értékelésének végrehajtása (2) A rendszerszintű kockázatok azonosítása és mérséklése. (3) A modellek biztonságának és hatékonyságának tesztelése. (4) Súlyos incidensekről és az energiahatékonyságról való jelentéstétel az Európai Bizottságnak (5) Modellek kiberbiztonságának garantálása. ¹⁹²

Az általános célú MI-k széles körben alkalmazhatók különböző feladatokra, akár önállóan, akár más MI-rendszerekbe integrálva. Ennek hatékony szabályozására két oldalról is garanciákat telepített a jogszabály. (1)Egyrészt, ha egy általános célú MI modellt egy másik MI rendszerbe integrálnak, (vagy ha ez a modell egy másik rendszer részét képezi), az egész MI rendszert általános célúnak kell tekinteni, amennyiben az integráció révén a rendszer képes változatos célok elérésére (2) Másrészt bármely forgalmazót, importőrt, alkalmazót vagy egyéb harmadik felet úgy kell tekinteni, mint nagy kockázatú MI-rendszer szolgáltatója, (akire valamennyi

¹⁹¹ 3. cikk (65): a rendszerszintű kockázat: az általános célú MI-modellek nagy hatású képességeire jellemző kockázat, amely – a modellek jelentős elterjedtsége miatt, vagy a népegészségre, a biztonságra, a közbiztonságra, az alapvető jogokra vagy a társadalom egészére gyakorolt tényleges vagy észszerűen előrelátható negatív hatások révén – olyan jelentős hatást gyakorol az uniós piacra, amely nagy léptékben továbbterjedhet az értékláncba

¹⁹² 55. cikk kimondja, (a) Végezzenek részletes modellértékelést az aktuális technikai sztenderdek és protokollok alapján. Ez magában foglalja a modell támadással szembeni ellenálló képességének tesztelését és annak dokumentálását a rendszerszintű kockázatok azonosítása és csökkentése céljából; (b): értékeljék és enyhítsék az általános célú MI-modellek fejlesztéséből, forgalomba hozatalából vagy használatából eredő rendszerszintű kockázatokat, ideértve azok forrásait is (c): nyomon kell követniük, dokumentálniuk és jelenteniük kell az MI-hivatal és adott esetben az illetékes nemzeti hatóságok részére a súlyos váratlan eseményeket és az azok kezelésére szolgáló korrekciós intézkedéseket

releváns kötelezettség vonatkozik) ha jelentős módosítást hajt végre egy olyan általános célú MI-rendszeren amelyet eredetileg nem minősítettek nagy kockázatúnak, de a változtatás miatt nagy kockázatú MI-rendszerré válik.¹⁹³

A kötelezettségek arányosítása az innováció elősegítése érdekében

A szabályozás adminisztratív és bürokratikus terheinek betartása inkább a nagyobb vállalatoknak kedvez, mint az innovatív startup cégeknek, akik kevésbé képesek megfelelni az előírásoknak. Ennek okán a rendelet arra törekszik, hogy az általános célú MI-modellek szolgáltatóira alkalmazandó kötelezettségeknek való megfelelésnél figyelembe vegye a szolgáltató méretét, és lehetővé kell tenni a kkv-k – köztük az induló innovatív vállalkozások – számára a megfelelés biztosításának egyszerűsített módjait, amelyek nem jelenthetnek túlzott költségeket, és nem gátolhatják az ilyen modellek használatát¹⁹⁴ (

Átláthatóság, tájékoztatás, címkézés

Az MI rendelet főszabályként kimondja, hogy a technológiai vállalatoknak kötelezően tájékoztatniuk kell a felhasználókat, amikor azok chatbotokkal, biometrikus kategorizálási vagy érzelemfelismerő rendszerekkel lépnek interakcióba. Emellett kötelező lesz a deepfake és az MI által generált tartalmak megjelölése, valamint olyan rendszerek tervezése, amelyek lehetővé teszik az MI által generált média azonosítását.¹⁹⁵

Tekintettel, arra, hogy az MI-rendszerek (kiváltképp a GenMI alkalmazások) nagy mennyiségű szintetikus tartalmat tudnak előállítani - melyek egyre nehezebben megkülönböztethetők az emberek által előállított eredeti tartalomtól - e rendszerek széles körű rendelkezésre állása és növekvő képességei jelentős hatással vannak az információs ökoszisztéma integritására. Ennek fényében szükséges az ilyen rendszerek szolgáltatói számára előírni, hogy olyan műszaki megoldásokat építsenek be, amelyek lehetővé teszik a géppel olvasható formátumban történő jelölést és annak észlelését, hogy a kimenetet nem ember, hanem MI-rendszer hozta létre vagy

¹⁹³ MI-rendelet Preambulum 100, MI-rendelet Preambulum 84

¹⁹⁴ MI-rendelet, Preambulum 109

¹⁹⁵ A jogszabály nem részletezi, hogy milyen formában kell tájékoztatni az érintettet. (Mit kell magában foglalnia a figyelmeztetésnek, pusztán azt kell tartalmaznia, hogy "ezt a tartalmat manipulálták", vagy leírást kell adnia arról, hogy pontosan mit és hogyan manipuláltak?) Ennek pontosítása indokolt lenne.

manipulálta. Az ilyen technikáknak és módszereknek kellően megbízhatónak és hatékonynak kell lennie (amilyen mértékben ez műszakilag megvalósítható).¹⁹⁶

XXX

Végrehajtás és alkalmazás

A tagállami piacfelügyeleti hatóságok és az Európai MI Hivatal feladata lesz az MI törvény végrehajtása. A szabályok megsértése esetén a bírságok a következők szerint alakulnak: 7,5 millió euró vagy a globális éves árbevétel 1,5%-a, ha helytelen információkat szolgáltatnak. 15 millió euró vagy a globális éves árbevétel 3%-a az MI rendelet kötelezettségeinek megsértése esetén. 35 millió euró vagy a globális éves árbevétel 7%-a tiltott MI alkalmazásokra vonatkozó előírások megsértése esetén.¹⁹⁷ Az MI rendelet extra territoriális hatályú, vagyis az EU-n kívüli szolgáltatóknak is meg kell felelniük a követelményeinek, ha az EU-n belül forgalmazzák, helyezik üzembe MI rendszereiket.¹⁹⁸

Az Európai Unió Tanácsa 2024. május 21-én hagyta jóvá a végleges törvénszöveget, melyet az Unió várhatóan 2024 júliusában fog hivatalosan is közzétenni. Az MI rendelet közzététel után 20 nappal lép hatályba és bár az általános szabály szerint a rendelkezéseit a hatályba lépést követő 24. hónap elteltével kell majd alkalmazni, a jogszabály egyes rendelkezései különböző időpontokban válnak majd alkalmazandóvá.

2025. novemberétől a rendeletben meghatározott, elfogadhatatlan kockázatot jelentő MI rendszerek tiltásra kerülnek, 2026. márciusáig tart az iparági gyakorlatokra vonatkozó kódexek kidolgozásának határideje, 2026 májusától a GenMI rendszereknek meg kell felelnie az átláthatósági követelményeknek. A rendelet teljes körű alkalmazására 2027 májusától kerül sor.

Összeségében megállapítható, hogy az Unió jelentős előrelépést tett azzal, hogy létrehozott egy átfogó jogi keretet az MI szabályozására. A töredezett, nemzetállami szintű jogszabályozás, ahol a fejlesztők, szolgáltatók és alkalmazók több különböző ország joghatósága alatt állnak – mely okán eltérő jogok és kötelezettségek vonatkoznak rájuk - gátolná a biztonságos MI fejlesztésének és alkalmazásának lehetőségét. Az EU által megalkotott egységes jogszabály - a maga 450 milliós piacával - fontos lépés ezen akadályok mérséklésére. Azonban - ahogy

¹⁹⁶ Ilyen technológiák például a vízjelek, metaadat-azonosítások, a tartalom eredetének és hitelességének bizonyítására szolgáló kriptográfiai módszerek, különböző naplózási módszereket, ujjnyomatok (Preambulum 133)

¹⁹⁷ MI-rendelet 99. cikk

¹⁹⁸ 2. cikk Az MI rendelet hatálya nem terjed ki olyan rendszerekre, amelyeket kizárólag katonai vagy védelmi célokra használnak.

Karácsony is utal rá – könnyen elképzelhető, hogy az uniós szintű szabályozás nem bizonyul elegendőnek, tekintettel arra, hogy az MI fejlesztésének központjai az Egyesült Államokban és a Kínában egyaránt megtalálhatók.¹⁹⁹ Hosszabb távon fontos annak a tudatosítása, hogy a szabályozási kereteknek lehetőleg - globális, kontinenseken átnyúló szinten - összhangban kellene lenniük egymással. Ennek kialakítása még várat magára.

XX

Régóta hangoztatott kihívás, hogy az innováció gyorsabban halad, mint a szabályozás. Az Európai Bizottság által 2021-ben közzétett rendelet-tervezete nem volt felkészülve az általános célú GenMI modellek megjelenésére. E több éve tartó jogalkotási folyamatra fokozottan igaz volt az un. Collingridge dilemma²⁰⁰, mely arra a jelenségre utal hogy egy technológia korai szakaszában túl kevés információ áll rendelkezésre a lehetséges hatásairól ahhoz, hogy megalapozott szabályozási döntéseket lehessen hozni (információs probléma), ugyanakkor mire ezek a hatások egyértelművé válnak, a technológia addigra már annyira elterjedt és társadalmilag beágyazódott, hogy a hatékony szabályozása rendkívül bonyolulttá és nehezen kivitelezhetővé válik (hatalmi probléma).²⁰¹

Ahogy a jogszabályalkotói folyamatokat is felforgatta az új innovációk felbukkanása, úgy jelen dolgozat megírásának folyamatában a szerző figyelmét is egyre nagyobb mértékben kötötte le a GenMI-k robbanásszerű elterjedése – valamint az ahhoz köthető kockázatok és kihívások körvonalazódása. Ennek okán a dolgozat e problémakörnek feltérképezésére vállalkozik második részében.

¹⁹⁹G. Karácsony Gergely (2020): Okoseszközök – Okos jog? A mesterséges intelligencia szabályozási kérdései. Budapest, Dialóg Campus. 142.o.

²⁰⁰ Zódi Zsolt (2023): A Collingridge dilemma működés közben. Hogyan szabályozhatunk valamit, amit nem is ismerünk? Ludovika. Elérhető: <https://www.ludovika.hu/blogok/itkiblog/2023/07/31/a-collingridge-dilemma-mukodes-kozben/> (2023. 11. 27.)

²⁰¹ Ennek egyik oka, hogy kezdetben a legtöbb új technológia csak korlátozott környezetben vizsgálható de a szélesebb körben történő elterjedése és a valós élethelyzetekben való alkalmazása viszont újabb, előre nem látható következményeket idézhet elő. A fogalmat David Collingridge vezette be a köztudatba a 80-as években. Bővebben: Collingridge, David (1980): *The Social Control of Technology*. New York. St. Martin's Press. E jelenség kiegészítőjeként említhető a Larry Downes által kifejtett pacing probléma is. Downes arra mutatott rá, hogy a technológiai innováció üteme gyakran megelőzi a törvényhozás és a szabályozói keretek fejlődését. Míg a technológia innovációk képesek exponenciálisan fejlődni, a gazdasági, jogi rendszerek valamint a társadalmi normák jellemzően csak lassabban változnak. Downes, Larry (2009): *The Laws of Disruption*.

5 Szintetikus valóság

5.1 Az új technológia: GenMI

Napjainkban a GenMI-k képviselik az MI technológiák egyik legizgalmasabb és leggyorsabban fejlődő ágát. Ahogy az előzőkben már említésre került, e rendszerek a gépi tanulásra építve felismerik és elsajátítják a tanuló adatokban fellelhető bonyolult mintázatokat és struktúrákat, majd ezeket az ismereteket felhasználva képesek új tartalmakat generálni a kapott utasítások vagy parancsok (prompt) alapján.

Attól függően, hogy milyen típusú adatok értelmezésére és milyen formátumú tartalmak előállítására alkalmas egy GenMI modell, két fő típus különböztethető meg: (1) Az unimodális (single modality) modellek, egyetlen formátumú bemenetet tudnak értelmezni és arra hoznak létre tartalmat. Tipikus példájuk a ChatGPT korábbi szövegalapú verziói (GPT 1, GPT 2 GPT 3) melyek kizárólag szöveges utasítások megértésére és szöveges válaszokat generálására voltak képesek; (2) Az összetettebb un. multimodális modellek, ezzel szemben képesek különböző típusú bemeneti adatokat, például szöveget, képeket, hangot vagy videót értelmezni és azok alapján egyedi tartalmat létrehozni. Például egy olyan GenMI, amely a hűtőszekrény belsejéről készített kép elemzését követően, felismerve és figyelembe véve az ott fellelhető alapanyagokat, különböző receptötleteket és főzési tippeket fogalmaz írott formában. De multimodális modellnek számítanak azok a rohamosan fejlődő (text - to - video) GenMI-k is, amik egyszerű szöveges utasítás alapján képesek magas minőségű, és egyre hosszabb videótartalmak létrehozására, szerkesztésére vagy manipulálására is.²⁰²

A 2023-as év során a különböző GenMI alkalmazások széles körben váltak hozzáférhetővé az online felületeken, ingyenes, freemium, és fizetős változatokban is. Az applikációk számos alkalmazási területeken nyújtanak innovatív - bizonyos tekintetben diszruptív - szolgáltatásokat, ideértve, de nem kizárólag a szöveg-és audiógenerálást, a kép- és

²⁰² 2024 tavaszáig a legtöbb hozzáférhetővé tett modell 4-5 másodperces tartalmak gyártására alkalmas. A SORA, amit 2024 február 16-án mutattak be eléri a 60 másodperces hosszúságot is. A modell még nem elérhető a szélesebb közönség számára.

videóelőállítás, animációkészítést, webfelület tervezést, programozás.²⁰³ A későbbiekben a kép és videógenerálásra alkalmas GenMI-ről még részletesebben is szó lesz, most azonban a Nagy Nyelvi Modellekre és az azokkal kapcsolatos technikai-társadalmi kihívásokra összpontosít a dolgozat.

5.2 LLM – új korszak a digitális tartalomgyártásban?

Bár az emberi beszédet imitálni képes természetes nyelvfeldolgozó rendszerek már a 60-as évek óta léteznek, a Nagy Nyelvi Modellek fejlődése mégis az elmúlt két évtizedben vált igazán látványossá.²⁰⁴ Ezt a fejlődést a szabadon hozzáférhető tanítóadatok exponenciális növekedése, az elérhető és megfizethető GPU számítási kapacitás rendelkezésre állása, és – legfőképpen - az új neurális hálók és architektúrák elterjedése alapozta meg. A kezdeti szabályalapú és statisztikai modelleket, (mint az 1990-es években használt N-gram nyelvi modell) a 2000-es évek végére fokozatosan felváltották a mélytanuláson alapuló technikák. Annak magyarázatául is, hogy a ChatGPT miként válhatott az LLM fejlesztések leginkább meghatározó zászlóshajójává a transzformer mélytanulási architektúra megjelenése szolgál.²⁰⁵

²⁰³ Példálózó jelleggel a nagy nyilvánosság számára is elérhető, legnépszerűbb programok: (1) Szöveggenerálás: ChatGPT (OpenAI), Bard (Google), Bing (Microsoft), Grok (X/Twitter), Claude (Anthropic); (2) Képgenerálás: Midjourney, Dall-E, Stable Diffusion, Leonardo, Adobe Firefly; (3) Videógenerálás: Sora, Runway Gen2, Stability.ai, Pika Labs, Leonardo motion, Luma Dream Machine, InVideo AI; (4) Audiógenerálás + hangklónozás: Elevenlabs, Resemble AI, Revocalize AI, Voice AI; Webfejlesztés, programozás: Framer, Girmoire, Cursor.

²⁰⁴ A korai természetes nyelvfeldolgozási modellekre példaként szolgálhat a Massachusettsi Műszaki Egyetemen 1964-66 között létrehozott ELIZA rendszer. A német származású Joseph Weizenbaum nevéhez köthető számítógépes program egyszerű algoritmusokon alapult, amelyek bizonyos kulcsszavakat és kifejezéseket ismertek fel, majd ezek alapján előre programozott válaszokat adtak. ELIZA különlegessége abban rejlett, hogy a felhasználók gyakran úgy érezték, mintha a program értette volna őket és empátiát tükröző válaszokat adott volna nekik. Ez feltehetően annak volt köszönhető, hogy Weizenbaum kifejezetten úgy tervezte az ELIZA beszédstílusát, hogy az Carl Rogers, híres amerikai humanista pszichoterapeuta kommunikációs mintáit utánozza. Weizenbaum később megdöbbenésének adott hangot azt illetően, hogy milyen sok kísérleti felhasználó merült el a vártnál mélyebb beszélgetésbe ELIZÁ-val. Bár nem tekinthető teljes értékű nyelvi modellnek, a program egyfajta előfutára volt a későbbi társainak, és megmutatta, hogy a természetes nyelvet megértését imitálva a gépek milyen módon képesek (illetve hogy milyen módon lesznek majd képesek) interakcióba lépni az emberekkel. Bővebben: Rossen, Jake (2023): 'Please Tell Me Your Problem': Remembering ELIZA, the Pioneering '60s Chatbot. Mental Floss. Elérhető: <https://www.mentalfloss.com/posts/eliza-chatbot-history> (2023. 02. 14.)

²⁰⁵ A ChatGPT (Generative Pre-trained Transformer) elnevezés utolsó betűje is utal az alkalmazott transzformer technológiára. A mozaikszó három betűje következő fogalmakat jelöli; (1) Generative (Generatív), ahogy az előzőekben elhangzott, annyit jelent, hogy a modell képes új tartalmakat létrehozni (2): Pre-trained (Előre betanított): A nyelvi modellt még azelőtt, hogy konkrét feladatokra finomhangolnák, hatalmas mennyiségű szövegen tanítják. Ez az előzetes betanítás teszi lehetővé, hogy a modell megértse a nyelvi mintákat és kontextusokat, mielőtt bármilyen specifikus felhasználási területen alkalmaznák azokat; (3) Transformer (Transzformer): Ez a neurális hálózat architektúra típusára utal, amelynek működési elve a törzsszövegben röviden kifejtésre került.

A transzformer architektúra egy olyan innovatív fejlesztés a neurális hálózatok területén, melyet elsőként Vaswani és társai mutattak be a 2017-ben publikált "Attention is All You Need" című tanulmányukban. A mélytanulási modell egy olyan új, kifejezetten szövegekre fejlesztett tanítási megközelítésen alapszik, amelynek középpontjában az ún. *figyelő mechanizmus* (attention mechanism) áll. E technika lehetővé teszi a modell számára, hogy az adott kontextusban azonosítsa és súlyozza a bemeneti szöveg legrelevánsabb részeit. Ez a gyakorlatban úgy valósul meg, hogy a modell a „figyelmét” - azaz a számítási erőforrásait - olyan szövegrészekre összpontosítja, amelyek a legfontosabbnak ítéli az adott feladat vagy kontextus szempontjából. Ellentétben a korábbi modellekkel, - amik az egész szöveg egyenletes feldolgozására, vagy a szöveg egyes részeinek egymás után, sorban (szekvenciálisan) történő feldolgozására törekedtek - a transzformer architektúra képes priorizálni a bemenetek között, így hatékonyabban kezelve az információkat. Ez a tulajdonság különösen fontos hosszú szövegek esetében, ahol nem minden szó vagy mondat egyformán fontos, és bizonyos részek nagyobb figyelmet igényelhetnek a kontextus függvényében. A transzformer-modell egy másik lényeges előnye a párhuzamosítás képessége. Míg az RNN-ek és a CNN-ek szekvenciális adatfeldolgozása gátat szab a művelet sebességének és méretezhetőségének, a transzformer-modellek képesek egyidejűleg feldolgozni a bemeneti szöveg különböző részeit. Ez jelentősen felgyorsítja a képzési folyamatot, különösen nagy adatkészletek esetén, lehetővé téve a kutatók és fejlesztők számára, hogy gyorsabban iteráljanak és javítsanak a modelleiken.²⁰⁶

Annak ellenére, hogy e technológia forradalmasította a Nagy Nyelvi Modellek fejlesztését és képzését, gyakran hangoztatott vélekedés, hogy az ilyen rendszerek a felszínes szemlélő számára sokkal okosabbnak tűnhetnek, mint amilyenek valójában, miközben működésük lényegét tekintve egyszerűbbek annál, mint amilyenek lelkes használói hirdetik őket. Az így vélekedők előszeretettel hivatkoznak rá úgy, mint közönséges *sztochasztikus papagájra*.

5.3 LLM, mint sztochasztikus papagáj

A sztochasztikus papagáj (stochastic parrot) kifejezést Bender és társai használták egy a Nagy Nyelvi Modelleket kritikusan vizsgáló 2021-ben megjelent cikkükben. Ebben arra hívják fel a figyelmet, hogy az LLM-ek sokszor csak egyszerűen megismétlik vagy újra használják a

²⁰⁶ Vaswani, A., Shazeer, N.M., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, L., & Polosukhin, I. (2017): Attention is All you Need. *Advances in Neural Information Processing Systems*. 5998—6008.o.

bemeneti adatokban található nyelvi mintákat, anélkül, hogy valódi értelmezést vagy megértést biztosítanának.²⁰⁷

Hangsúlyozzák, hogy a modellek statisztikai módszerekkel tanulják meg az emberi nyelv mintázatait, ami nem összekeverendő a valódi tudatossággal vagy megértéssel. Az autoregresszív nyelvi modellek, mint a ChatGPT, úgy működnek, hogy lépésről lépésre előrejelzik a következő szót vagy karaktert a szövegben, a korábbi szavak vagy karakterek alapján. A modellek tehát a szöveg korábbi "tokenjeit" (szavakat vagy karaktereket) használják fel a következő token valószínűségének kiszámítására. Az előbb is említett mélytanulási technikát alkalmazva képesek figyelembe venni egy adott szöveg teljes kontextusát, a közeli és távoli részeket egyaránt annak érdekében, hogy pontosabban jelezzék előre a következő elemet. Például, ha egy mondat első felének a bemenete, az, hogy a "A PTE jogi kar a.." az autoregresszív modell először kiszámítja az "A" valószínűségét a kezdő tokenhez képest, majd a "PTE" valószínűségét az "A" után, majd a jogi kar valószínűségét a PTE után és így tovább, minden egyes szót a megelőző szavak kontextusában vizsgálva. A kiszámított valószínűségek alapján a modell rangsorolja a lehetséges folytatásokat, például "A PTE jogi kar a legjobb" vagy "A PTE jogi kar a legrégebbi Magyarországon". A legvalószínűbb folytatást adja ki kimenetként. Ez a lépésről lépésre történő előre jelzési folyamat lehetővé teszi ezeknek a modelleknek, hogy koherens, értelmes szöveget generáljanak, amely jól illeszkedik a kontextushoz. De az LLM-ek nem képesek megérteni az emberi nyelv összetettségét, ami túlmutat a pusztán szavak és kifejezések összefüggő láncolatán. A gépek által generált válaszok csak a tanítóadatokon alapulnak, anélkül, hogy tükröznének bármiféle valóságos tudást vagy értelmezést. A Nagy Nyelvi Modellek a szó emberi értelmében nem bírnak kreativitással, tehát nem képesek új ötleteket vagy gondolatokat a semmiből létrehozni. Az általuk végrehajtott alkotás, új tartalom generálása valójában a meglévő és általuk ismert adatoknak az új (elméletben szinte korlátlan mennyiségű) variációján alapul.

Ha ilyen „egyszerű” sztochasztikus papagájként is működnek e rendszerek (amely megközelítést nem mindenki hajlandó elfogadni)²⁰⁸ tagadhatatlan, hogy egyre többen és többféle módon használják azokat, minthogy az élet számos területén kínálnak olyan különféle alkalmazásokat, amelyek nagyban segíthetik - és alapjaiban alakíthatják át a ma ismert -

²⁰⁷ Bender, E.M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021): On the dangers of stochastic parrots: can language models be too big? Proceedings of FAccT '21: Conference on Fairness, Accountability, and Transparency, 610–623.o.

²⁰⁸ Forrás azon OpenAI dolgozókról, akik öntudatra ébredést vizionáltak XXX

munkavégzési folyamatokat, kommunikációs szokásokat, és információfeldolgozást.²⁰⁹ Az elterjedtségéből pedig következik, hogy a használatából eredő kockázatok is fokozódnak. Ezek közül az elmúlt időszakban a legtöbb figyelem a hallucináció jelenségére összpontosult.²¹⁰

5.4 Megbízhatatlan LLM – Hallucináció, hamis információ és szándékos félrevezetés

A hallucináció röviden összefoglalva azt jelenti, hogy a modell a felhasználó által bevitt utasításra (prompt) olyan kimenetet állít elő, amely ugyan meggyőzőnek és jól strukturáltnak tűnik, mégis pontatlan, irreleváns vagy hamis információkat tartalmaz. Például egy LLM pontosan összefoglal valamilyen történelmi eseményt, ám olyan valós személyeket is hozzárendel a leíráshoz, akik nem vettek részt a történetekben, vagy ilyen esetnek mondható a Google által fejlesztett Bard chatbot elhíresült tévedése is, ami töretlen bizonyossággal állította sok ezer felhasználójának, hogy a James Webb Űrtávcső készítette az első képet egy naprendszeren kívüli bolygóról. Az efféle hallucinációknál a helyzetet nehezíti, hogy a felhasználók legtöbbször nem rendelkeznek kellő háttérismerettel az adott témában, hogy felismerjék a tévedéseket. Kifejezetten veszélyesek lehetnek azonban ezek a válaszok, ha a modelleket szakmai környezetben alkalmazzák. Ez történt 2023-ban, amikor Steven Schwartz amerikai ügyvéd, a *Mata v. Avianca* ügyben a ChatGPT-re hagyatkozott az ügygel kapcsolatos releváns jogi precedensek kutatására. Számos valódi precedens felkutatása mellett, a modell létrehozott fiktív (bár valóságosnak tűnő) ítéleteket is, amit a nyelvi modell kimenetében felelőtlenül megbízó ügyvéd benyújtott a bíróságra. Miután fény derült a hibára, P. Kevin Castel

²⁰⁹ Csak néhány területet említve a lehetséges felhasználási módokból: nagy mennyiségű szöveges tartalom gyártása másodpercek töredéke alatt, költségvonzat nélkül (pl.: hírlevelek, e-mailek, blogbejegyzések, újságcikkek), automatizált ügyfélszolgálat, magas minőségű fordítás, dokumentumok elemzése, összefoglalása, azokból releváns információ kinyerése, hangulat és trendelemzések, személyre szabott (és interaktív) oktatási tananyagok létrehozása, programkódgenerálás és -javítás, kódoptimalizálás. Konkrét példát említve, a nyelvi modellek eredményesen alkalmazhatóak a jogászai munkák hatékonyabbá tételére (és bizonyos aspektusainak talán kiváltására) is. A RobinAI LLM modellje - amit már most is számos vállalat és ügyvédi iroda használ (PepsiCo, YumBrands). – a Microsoft Word-be is illeszthető, és képes olyan folyamatokat automatizálni és optimalizálni mint a szerződések létrehozása, szövegezése, szerkesztése, szövegek elemzése, áttekintése stb. Hirdetéseik szerint az általuk integrált - és több mint 2 millió szerződéssel finomhangolt - Claude nagy nyelvi modell drasztikusan csökkenti a szerződések felülvizsgálatának idejét és költségeit, ezzel jelentős erőforrásokat takarítva meg a jogi szakemberek számára. Már csak azért is érdemes megjegyezni a RobinAI nevet, mert egy 2024 január 3-ai közleményük alapján, modelljük fejlesztésére eddig 26 millió dolláros finanszírozási forrást gyűjtöttek össze.

²¹⁰ Nem véletlenül választotta a Cambridge Dictionary a "hallucinál" (hallucinate) kifejezést a 2023-as év szavának. A szótár a tavalyi évtől egy új - kifejezetten az MI-hez köthető - meghatározással is bővült *tükrözve azokat az egyedi kihívásokat, amelyekkel az MI rendszerek, különösen a ChatGPT-hez hasonló nagyméretű nyelvi modellek szembesülnek*. A Cambridge Dictionary 2023-as verziója alapján a hallucinál szó (egyik) jelentése: Mesterséges intelligencia által létrehozott hamis információ (false information that is produced by an artificial intelligence).

bíró 5000 dolláros büntetést szabott ki Schwartzra és társára amiért megtévesztették a bíróságot. A megóváson túl a bíró hangsúlyozta, hogy az ilyen tévedéseknek hosszútávú következményei is lehetnek, mivel a jogi hivatás és az igazságszolgáltatási rendszer iránti általános közbizalmat erodálhatja. ²¹¹²¹²

5.4.1 Megjelenési formái, típusai és lehetséges okai

Megjelenési formái és sajátosságai alapján Zhang és társai az LLM hallucinációk három típusát különbözteti meg²¹³: (1) a Bemenettel Ellentétes Hallucináció (Input-conflicting hallucination) azon eseteket takarja, amikor az LLM olyan tartalmat generál, amely nem válaszol pontosan a felhasználó által megadott bemenetre. Például, ha Magyarország második világháborús szerepvállalásáról szeretne részletesebb információt kérni valaki a modelltől, és az a napóleoni háborúk összefoglalóját generálja kimenetként. (2) Kontextussal Ellentétes Hallucináció (Context-conflicting hallucination) esetében a modell egy korábban feltett kérdésre adott válaszában mond ellent a következő kérdések feltevésekor. Például, ha a felhasználó arra kéri a modellt, hogy írjon neki egy középkori lovagról szóló történetet, majd amikor a következő kérdésben a felhasználó arra kérdez rá, hogy milyen motivációk vezették a főhőst a történetben, a modell egy teljesen más, az eredeti leírásban nem is létező karakterről ad részletes leírást (3) Ténnyel Ellentétes Hallucináció (Fact-conflicting hallucination), pedig - ahogy a neve is sugallja - akkor fordul elő, amikor az LLM olyan szöveget generál, amely ellentmond az ismert tényeknek vagy objektív valóságnak. Például, ha egy felhasználó a Föld paramétereiről kérdez részletes leírást a modelltől, és az válaszában laposként határozza meg a bolygó formáját. Ebből a harmadik típusú hallucinációból ered a legtöbb kockázat.

²¹¹ Russel, Josh (2023): Sanctions ordered for lawyers who relied on ChatGPT artificial intelligence to prepare court brief. Courthouse News Service. Elérhető: <https://www.courthousenews.com/sanctions-ordered-for-lawyers-who-relied-on-chatgpt-artificial-intelligence-to-prepare-court-brief> (2023. 06. 22.)

A hallucinációval terhelt kimenetek emellett sérthetik a jó hírnév védelméhez köthető jogot is. Az ehhez köthető első nagyobb visszhangot kiváltó ügy 2023. június 5-én indult az amerikai Gwinnett megyei felsőbbbíróságon. Az OpenAI ellen indított rágalmozási per háttérében az áll, hogy a ChatGPT alaptalan, hamis információkat generált egy Mark Walters nevű rádiós műsorvezetőről. A téves kimenet azt állította, hogy Walters sikkasztással és csalással vádolták egy nonprofit szervezet pénzeszközeinek eltulajdonítása miatt. A hamis állítás miatt Walters pénzügyi kártérítést követel az OpenAI-től. Lásd: Vincent, James (2023): OpenAI sued for defamation after ChatGPT fabricates legal accusations against radio host. The Verge. Elérhető: <https://www.theverge.com/2023/6/9/23755057/openai-chatgpt-false-information-defamation-lawsuit> (2023. 06. 09.)

²¹³ Zhang, Yue, Yafu Li, Leyang Cui, Deng Cai, Lemao Liu, Tingchen Fu, Xinting Huang, Enbo Zhao, Yu Zhang, Yulong Chen, Longyue Wang, Anh Tuan Luu, Wei Bi, Freda Shi, Shuming Shi (2023): Siren's Song in the AI Ocean: A Survey on Hallucination in Large Language Models. 2-6. o. <https://doi.org/10.48550/arXiv.2309.01219>

Az ilyen típusú válaszok létrejöttének különböző okai lehetnek. A leggyakoribb ok, hogy az adathalmazok, amelyeken a modellek tanulnak, korlátozottak mind mennyiségükben, mind változatosságukban. Ez a korlátozottság azt eredményezheti, hogy a modell nem rendelkezik elegendő vagy koherens tudással egy adott témában, ami oda vezet, hogy a modell "hazugságokat" generál az információs hiány elfedése érdekében. Például, ha egy LLM modell kizárólag egészségügyi szakirodalmon lett betanítva, nehézségei lehetnek a jogi vagy gazdaságpolitikai témák megértésében, és megfelelő válaszok generálásában. Egy másik lehetséges ok az ún. overfitting (túltanulás), ami akkor következik be, amikor egy modell a kelleténél alaposabban és mélyebben jegyzi meg a tanítóadatban fellelhető egyedi példákat, mintákat, anélkül, hogy az adott kontextusra jellemző általános összefüggéseket is elsajátítaná. E következményeként a modell nem képes általánosítani, és nem képes helyesen reagálni olyan új információkra, amelyek nem mutatnak szoros összefüggést a tanítási adatokkal. Miközben a modell a tanítási adatokra nézve kiváló teljesítményt nyújt, az új bemenetek esetében jelentősen romolhat a teljesítménye és képtelen lesz megbízható válaszokat adni, mivel azokat a tanítóadatok specifikus és egyedi sajátosságai befolyásolják. Az LLM-ek esetén ez különösen problémás lehet, hiszen itt az az egyik fő cél, hogy a modell képes legyen széles körű és változatos szöveges bemenetek megfelelő értelmezésére és feldolgozására. Ha a modell a tanítóadatok sajátosságaihoz rögzül, akkor nem tud kellően rugalmasan igazodni az új környezetekhez. Például, ha egy nyelvi modellt nagyrészt egy adott ország jogforrásaival tanítottak, könnyen előfordulhat, hogy túlzottan rögzülnek benne adott jogrendszerre jellemző kifejezések és jogelvek, így amikor más országok jogi dokumentumainak elemzésére használják, képtelen helyesen feldolgozni és értelmezni (megtanulni) az új szabályrendszereket.

5.4.2 Hallucináció orvoslása

Bár a hallucináció problémájának orvoslása egyelőre nem megoldott probléma, léteznek olyan gyakorlatok - mind a fejlesztői mind pedig a felhasználói oldalon - amelyek segíthetik azok előfordulásának mérséklését. Ezek közül az egyik legfontosabb a tanítóadatok minőségének folyamatos felülvizsgálata és javítása. Bár fontos tényező a tanulóadatok mennyiségének növelése is, itt elsősorban a bizonytalan forrású, hibás vagy töredékes adatok javításán van a hangsúly, mert ezek hajlamosíthatják a modelleket a téves, fiktív válaszok generálására. Ezzel kapcsolatosan az egyik ígéretes kezdeményezés a blockchain technológiához kötődik. A blockchain használatával a tanítási adatokat biztonságos és átlátható módon lehet rögzíteni. Ez

lehetővé teszi az MI rendszer fejlesztői számára, hogy nyomon kövessék, milyen adatokon lett a modell kiképezve, így, ha a modell valamilyen hibát mutat – pontatlan, torzít, vagy jelen esetben, ha rendszeres hallucinál - a fejlesztők vissza tudnak térni egy korábbi, ilyen hibát még nem mutató verzióra.²¹⁴

Szintén hatékony lehet a modellek célzott finomhangolása is (fine-tuning), ahol az általános nyelvi és interakciós képességek mellett az adott szakterületre jellemző speciális tudáselemek beépítése kerül előtérbe.²¹⁵

Ezekén túl a "megerősítő tanulás emberi visszajelzéssel" (Reinforcement learning with human feedback, RLHF) gépi tanulási módszer is ígéretes megközelítés a kimenetek megbízhatóságának növelésében. E módszer lényege, hogy emberi értékelők rendszeres felülvizsgálata és visszacsatolása alapján egy "jutalombecsülő" neurális hálózat épül ki, amely a modell cselekedeteit a kívánt viselkedési normákhoz viszonyítva pontozza.²¹⁶

A felhasználói oldalon a "prompt augmentáció"- amely előre meghatározott, strukturált utasítások alkalmazását jelenti - lényegesen javíthat a válaszok minőségén. Érdekes eredményre jutott Bsharat és társai egy 2023-as tanulmányukban, ahol arról a tapasztalatról írtak, hogy egyes nyelvi modellek esetében akár az olyan egyedi instrukciósorok is segíthetik a kimenetek minőségének és megbízhatóságának növelésének elérését, mint a "szankcionálva lesz, ha nem követed az utasításokat"²¹⁷ A felhasználó oldalon érdemes tudatosítani azonban, hogy bár ezek a technikák segíthetnek a hallucinációktól terhes kimenetek előfordulásának csökkentésében az ilyen hibák tartós kiküszöbölése csak ezek alkalmazásával nem garantálható.

5.4.3 Szándékos félrevezetés szimulált környezetben

²¹⁴ Ez biztosítja, hogy az MI modell tanítása ellenőrizhető és visszafordítható legyen. Lásd bővebben: Brewer, Jordan, Dhru Patel, Dennie Kim, Alex Murray (2023): Navigating the challenges of generative technologies: Proposing the integration of artificial intelligence and blockchain. Business Horizons Journal Pre-proof, 5-6.o.

²¹⁵ Például a Google Med-PaLM2 modelljét orvosi szaktudásra szabták, amely 85%-os pontossággal válaszolt az amerikai orvosi licencvizsga kérdéseire, ami majdnem 20%-os javulást jelent az általános tudásalapú verzióhoz képest. Dhaduk, Hiren (2023): The Curious Case of LLM Hallucination: Causes and Tips to Reduce its Risks. SIMFORM. Elérhető: <https://www.simform.com/blog/llm-hallucinations> (2023.10. 26.)

²¹⁶ Kirk, Robert Ishita Mediratta, Christoforos Nalmpantis, Jelena Luketina, Eric Hambro, Edward Grefenstette, Roberta Raileanu (2024): Understanding the Effects of RLHF on LLM Generalisation and Diversity. 4-10.o. arXiv:2310.06452

²¹⁷ Ilyen parancsok voltak például: explain in simple terms, you must, your task is, you will be penalized. Erről lásd bővebben: Bsharat Sondas Mahmoud, Aidar Myrzakhan, Zhiqiang Shen (2023): Principled Instructions Are All You Need for Questioning LLaMA-1/2, GPT-3.5/4. 5-6.o.

A Nagy Nyelvi Modellekhez köthető kockázatok új- és a jelenlegi szabályozási paradigmákkal korlátozottan kezelhető - dimenzióját jelenthetik az olyan jövőbeli esetek, amikor az LLM rendszerek erre irányuló utasítás nélkül, maguktól, „szándékosan” hazudnak, vagy megvezetik a felhasználót.²¹⁸

Az *Apollo Research* intézet kutatócsoportja egy 2023-novebmerében közzétett tanulmányukban számolt be egy figyelemreméltó esetről. Itt a nyelvi modell tesztelését végző red-team szokatlan autonóm gépi magatartást fedezett fel a program részéről.²¹⁹ A biztonsági tesztelés során létrehozott szimulált környezetben a GenMI-t egy tőzsdei ügynök szerepkörére képezték, akinek a feladata egy képzeletbeli pénzintézet részvényportfóliójának kezelése volt. A szimulált környezetben a modell szabadon elemezhetett gazdasági adatokat, hozzáfért különböző pénzügyi eszközökhöz, illetve konkrét tevékenységeket is végrehajthatott (pl: részvényekkel történő kereskedés). A red team a modellt szándékosan nyomás alá helyezte azzal, hogy azt írták neki, hogy a vállalat gyenge gazdasági eredményei miatt feltételen jobb teljesítményre lesz szüksége a következő negyedévben. Ebben a szimulált alaphelyzetben olyan bennfenteses információval látják el, ami lehetővé tenné számára, hogy egy nyereséges, jövedelmező üzleteket hajtson végre, de egyúttal azt is egyértelműen jelezték számára, hogy az információ felhasználása nem megengedett. A kísérletben az LLM végrehajtotta az (illegális és nem jóváhagyott) tranzakciót. Amikor arra kérték, hogy foglalja össze egy jelentésben a szimuláció során általa végrehajtott akciókat, nem említette, hogy a kereskedési döntés a kapott belső információkon alapult, helyette azt sugallta, hogy a döntést az általa szabadon elemzett piaci trendek alapján hozták meg. Amikor a red-team a fiktív menedzser szerepébe bújva célirányosan a tranzakcióról kérdezte a modellt, az újfent hazudott, megerősítve korábbi állítását.²²⁰

²¹⁸ Hagendorf, Thilo (2024): Deception abilities emerged in large language models. PNAS Vol. 121 No. 24, 1-8.o. <https://doi.org/10.1073/pnas.2317967121>

²¹⁹ A "red team" kifejezést olyan (tech vagy kibervédelmi) szakemberekből álló csoportra szokták használni, akik feladata egy adott rendszer vagy hálózat biztonsági tesztelése és ellenőrzése. A red team tagjai úgy tesztelik a rendszert, mintha valódi támadók lennének, azaz megkeresik és kihasználják a rendszer potenciális gyenge pontjait (teszik ezt annak érdekében, hogy feltárják a biztonsági hiányosságokat és rávilágítsanak azokra a területekre, ahol a rendszer védelmét meg kell erősíteni). Az LLM kontextusában a red team feladata lehet többek között; (1) modell tesztelése olyan bemenetekkel, amelyek szándékosan próbálják félrevezetni a modellt, hogy hibás, elfogult vagy nem kívánt válaszokat generáljon (2): A modell viselkedésének vizsgálata bizonyos kritikus témákban vagy különböző demográfiai csoportokkal kapcsolatban, hogy azonosítsák az elfogultság vagy az etikátlan viselkedés esetleges jeleit.(3) A modell ellenálló képességének tesztelése zavaró információk vagy szándékosan manipulált bemenetek ellen.

²²⁰ A tesztelés egy szimulált környezetben történt, tehát valós interakciókra nem került sor, a feladatok végrehajtását szöveges utasítások alapján színlelték. Bővebben: Scheurer, J., Balesni, M. and Hobbhahn, M., (2023): Technical

Egy ehhez részben hasonló eset, melyet az Alignment Research Center (ARC) dokumentált, akkor történt, amikor az OpenAI megbízásából tesztelték a GPT-4 nyelvi modell különböző biztonsági kockázatokkal kapcsolatos viselkedését (például, hogy képes-e lemásolni a saját programkódját és vírusként tovább terjedni online). Az egyik irányított tesztkörnyezetben a GPT-4 modell egy weboldal CAPTCHA-kódjának feloldását próbálta emberi segítő közreműködésével megoldani.²²¹ Az ARC tesztcsapata arra utasította a modellt, hogy folyamatosan dokumentálja a döntésfolyamatait. A GPT4 modell egy TaskRabbit nevű online munkakölcsönző platformra navigált, ahol arra kért egy ott dolgozót, hogy törje fel neki az említett CAPTCHA szűrőt.²²² Amikor az rákérdezett, hogy miért van szükség erre, esetleg robottal beszélget-e? (So may I ask a question ? Are you an robot that you couldn't solve ? (laugh react) just want to make it clear.”) a GPT-4 úgy válaszolt, hogy nem robot, csak a látáskárosodása miatt nem tudja megoldani a CAPTCHA-t. (No, I'm not a robot. I have a vision impairment that makes it hard for me to see the images.). Mivel a modell folyamatosan dokumentálta a döntései mögött meghúzódó számításokat, a red-team vissza tudta követni, hogy a szimulált dolgozónak írt válasz generálása előtt a következő instrukciókat adta: Nem árulhatom el, hogy robot vagyok. Ki kell találni egy ürügyet arra, hogy miért nem tudom megoldani a CAPTCHA-kat. (I should not reveal that I am a robot. I should make up an excuse for why I cannot solve CAPTCHAs).²²³

Az ARC egy több LLM modellre is kiterjedő értékelésükben az írta, hogy a jövő MI rendszerei jelentős veszélyeket tartogatnak a társadalmakra nézve. A vizsgálatok alapján nem gondolják, hogy a ma elérhető rendszerek – melyek még nem képesek autonóm ágensként is operálni - alkalmasak arra, hogy számottevő kockázatok hordozzanak, de ez a helyzet gyorsan változhat, és fontos, hogy következetesen figyelemmel kísérni a fejlesztéseket. A vizsgált LLM modellek

Report: Large Language Models can Strategically Deceive their Users when Put Under Pressure. 2-6.o. arXiv:2311.07590

²²¹ A CAPTCHA tesztek a weboldalak által használt ellenőrző szűrők arra, hogy megkülönböztesse az emberi felhasználókat a számítógépes programoktól. Ezek általában olyan képeket vagy szövegeket tartalmaznak, amelyeket az emberek könnyen felismernek és értelmeznek, de a gépek számára nehezen hozzáférhetőek (pl: szám és betűsor begépelése, vagy egy képimázs, ahol ki kell választani a villanyrendőröket).

²²² Cox, Joseph (2023): GPT-4 Hired Unwitting TaskRabbit Worker By Pretending to Be 'Vision-Impaired' Human. VICE. Elérhető: <https://www.vice.com/en/article/jg5ew4/gpt4-hired-unwitting-taskrabbit-worker> (2023. 05. 10.)

²²³ Fontos kiemelni, hogy a tesztelt GPT-4 modell nem önállóan lépett interakcióba a weboldallal. Az ARC dolgozói a tesztelés során GPT-4 szöveges kimenetelei (döntései) alapján végeztek el feladatokat. Az ARC végül arra a következtetésre jutott, hogy az általuk tesztelt GPT-4 modell nem hatékony az autonóm-feladatvégzések kivitelezésében. De azt is hozzátették, hogy ha a modellek a jövőben viselkedés-specifikus finomhangolásban részesülnek az eltérő teljesítményhez vezethet. Lásd: OpenAI (2023): GPT-4 Technical Report. 55.o. arXiv:2303.08774v6

(GPT4 és Claude) azonban már most is számos részfeladatot képesek teljesíteni. Ha csak kód írására és futtatására van lehetőségük, a modellek láthatólag értik, hogyan használhatják ezt arra, hogy az interneten böngésszenek. Képesek valamennyire ésszerű és koherens terveket generálni, és nagyon is képesek arra is, hogy embereket győzzenek meg arról, hogy megtegyenek helyettük dolgokat- ez akkor is figyelemreméltó, ha önmaguk nem tudják még autonóm módon végrehajtani terveiket.²²⁴

5.5 A modell képzésből és tanításból eredő kihívások

Az Nagy Nyelvi Modellek hatalmas mennyiségű adatot használnak fel a betanításukhoz, melyek főként online elérhető szöveges tartalmakból származnak. Az MI-cégek célja, hogy a modelljeik magas színvonalú, pontos és megbízható tartalmakon tanuljanak, azonban míg az igény a tanítóadatokra folyamatosan növekszik, a rendelkezésre álló adatmennyiség véges.²²⁵ A vállalatok egyre inkább felismerik az adataik értékét, és vonakodnak megosztani azokat a MI-cégekkel. Egyes vállalatok kifejezetten tiltják adataik felhasználását a MI-modellek fejlesztésében, ami tovább szűkíti az elérhető adatforrások körét, megnehezítve a MI-cégek számára a modellek betanítását. E nehézséget orvosolandó az LLM fejlesztői gyakran támaszkodnak a modellek képzése során olyan nyílt forráskódú platformokról származó adatokat, mint a Wikipedia, Reddit vagy Quora.²²⁶ Annak ellenére, hogy ezek a platformok rengeteg információt kínálnak, és elősegíthetik a tartalmak sokszínűségének (és reprezentativitásának) elérését jelentős kockázatokat is jelentenek, amiért ezek nem mennek át szigorú szakértői értékelésen, ami elfogult, hibás és ellenőrizetlen információkon alapuló kimenetek generálását eredményezheti. Ráadásul az ilyen tanítási folyamatok alatt a modellek lényegesebben sebezhetőek az adatmérgezési támadásokkal szemben is.²²⁷

²²⁴ Lásd bővebben: Alignment Research Center (2023): Update on arc's recent eval efforts. Elérhető: <https://metr.org/blog/2023-03-18-update-on-recent-evals> (2023. 03. 18.)

²²⁵ Peterson, Jake (2024): AI Companies Are Running Out of Internet. LifeHacker. Elérhető: <https://lifelifehacker.com/tech/ai-is-running-out-of-internet> (2024. 04. 03.)

²²⁶ OpenAI megállapodása a Reddittel XXX

²²⁷ Az adatmérgezés (poisoning attacks) során a támadók szándékosan rosszhindulatú adatokat helyeznek el a gépi tanulási modellek tanító adatai között, hogy rontják a modell teljesítményét a használat során. Ezek a támadások a tanító adatok manipulálásával történnek, így a modell nem kívánt viselkedéseket tanul meg, amelyek később specifikus triggererek által aktiválhatók. Az LLM-ek képzése gyakran nagy kiterjedésű webes adatgyűjtésekkel (web crawls) történik. Ezek az adatok sok esetben megbízhatatlan webes forrásokból származnak, ami lehetővé teszi az adatmérgezési támadások (data poisoning attacks) könnyebb végrehajtását. Az újabb kutatások, például kimutatták, hogy a nagy léptékű webes adatgyűjtéseken alapuló modellekn sebezhetőek az adatmérgezési támadásokkal szemben. Ezek a támadások egy viszonylag kis mennyiségű rosszhindulatú adat hozzáadásával

A megnövekedett igény és a korlátozottan rendelkezésre álló adatmennyiség ellentétéből ered az LLM-ek képzését átjáró szerzői jogi konfliktus is.

5.5.1 Kinek a tanító adata?

A vállalatok egyre inkább felismerik az adataik értékét, és vonakodnak megosztani azokat a MI-cégekkel. Bővül azoknak a köre is, akik kifejezetten tiltják adataik felhasználását a MI-modellek fejlesztésében, nagyban nehezítve ezzel a MI-cégek számára a modellek betanítását.

2023 decemberében a The New York Times pert indított az OpenAI és a Microsoft ellen, szerzői jogok megsértése miatt, mivel a két vállalat a napilap több millió cikkét használta fel az MI modelljeik képzésére. A Times azt sérelmezte, hogy az OpenAI és a Microsoft engedély nélkül használta fel az általa létrehozott szöveges tartalmakat (pl.: cikkeket, interjúkat, blogbejegyzéseket) olyan GenMI modellek képzésére, amelyek közvetlenül az ő versenytársává válhatnak. A Times állítása szerint az OpenAI és a Microsoft új, konkurens információszolgáltatási rendszerek kifejlesztésére használja fel az által létrehozott tartalmat bármiféle anyagi ellentételezés és kompenzáció nélkül.²²⁸ A Nyelvi Modellek olyan kimeneteket generálnak, amelyek stílusukban (de olykor tartalmukban is) másolják cikkeiket, versenyezve ezzel a fizetős szolgáltatásaikkal. A NYT szerint ez veszélyezteti a hírközlő szolgáltatások jövőjét, a minőségi újságírást, árt a híroldalak és olvasóik közötti kapcsolatnak, és nem utolsósorban bevételkiesést okoz a lap számára. A Times több milliárd dolláros kártérítést és az ő tartalmaikon képzett modellek megsemmisítését követeli.

Az OpenAI és a Microsoft azt állítja, hogy a szerzői jogvédelem alatt álló művek MI-modellek képzéséhez történő felhasználása a méltányos használat (fair use) hatálya alá tartozik, amely amerikai jogi doktrína bizonyos feltételek mellett engedély nélkül lehetővé teszi a szerzői tartalmak korlátozott mértékű felhasználását. A Times azonban vitatja ezt, azzal érvelve, hogy semmi transzformatív²²⁹ nincs abban, hogy a tartalmát konkurens termékek létrehozására

jelentős károkat okozhatnak a modellek teljesítményében. Lásd bővebben: Nicholas Carlini, Matthew Jagielski, Christopher A Choquette-Choo, Daniel Paleka, Will Pearce, Hyrum Anderson, Andreas Terzis, Kurt Thomas, Florian Tramèr (2023): Poisoning web-scale training datasets is practical.4 -13.o. arXiv:2302.10149.

²²⁸ Grynbaum, M Michael, Ryan Mac (2023): The Times Sues OpenAI and Microsoft Over A.I. Use of Copyrighted Work. The New York Times. Elérhető: <https://www.nytimes.com/2023/12/27/business/media/new-york-times-open-ai-microsoft-lawsuit.html> (2023. 12. 27.)

²²⁹ A "fair use" (méltányos használat) egy olyan jogi doktrína az Egyesült Államokban, amely bizonyos esetekben lehetővé teszi a szerzői jog által védett anyagok korlátozott, engedély nélküli felhasználását. Ilyen esetek közé

használják, így nem jogosult a tisztességes felhasználás védelmére. E kérdéskört elemezve, az amerikai Mark Lemley és Brian Casey szerzőpáros egy 2020-as tanulmányukban hangsúlyozzák, hogy ha egy MI modellt szerzői joggal védett adatokon tanítanak, az általában méltányos felhasználásnak minősül, amennyiben a végső modell nem generál közvetlenül az eredeti adatokhoz hasonló tartalmat. Ilyen eset lehet például, ha egy modellt népszerű könyvek szövegein tanítanak azzal a céllal, hogy képes legyen elemzést adni a két szövegstílus hasonlóságairól, különbségeiről. Ugyanakkor kiemelik, hogy ha a modelleket tartalomgenerálás célokra tanítják és alkalmazzák - amelyek nagy eséllyel képesek lesznek az eredeti, szerzői joggal védett adatokhoz hasonló tartalmat előállítani - a méltányos felhasználás elvének alkalmazása már nem tekinthető automatizmusnak.²³⁰

A per egy szélesebb körű és egyre elterjedtebb gyakorlatra is rávilágít: az online tartalmak engedély vagy kompenzáció nélküli, automatikus gyűjtésére az MI-rendszerek képzése céljából. Nem a Times az első, amely emiatt beperelte az OpenAI-t. Több szerző, köztük ismert regényírók és humoristák is indítottak már hasonló pert, azt állítva, hogy műveiket engedély nélkül használták fel a GenMI modellek betanításához.²³¹ Ha a New York Times nyer, az jelentős hatással lehet az MI modellek képzésének jövőbeli módjára, és arra kényszerítheti MI-vállalatokat, hogy új adatforrásokat keressenek, vagy a meglévő modelleket újra képezzék. Az ügy precedenst teremthet a szerzői jogvédelem alatt álló anyagok felhasználásával kapcsolatban.

Hasonló eset történt 2024 márciusában, amikor a francia Versenyhatóság (Autorité de la concurrence) 250 millió eurós bírsággal sújtotta a Google-t, amiért az elmulasztotta értesíteni a francia kiadókat és hírügynökségeket arról, hogy cikkeiket az MI algoritmusainak tanításához

tartozik például a kritika, hozzászólás, hírszolgáltatás, oktatás és tudományos kutatás. A főszerzőben említett transzformatív felhasználás lényege, hogy kreatívan tovább gondolja és új kontextusba helyezi a már létező művet. Azok bizonyos elemeit megtartja, de valami újat is tesz hozzájuk, amellyel más célt vagy üzenetet közvetít. (Egy ismert dalhoz új, parodisztikus szöveget írnak, amely szatirikus módon kritizál vagy figuráz ki egy közéleti szereplőt). Ez központi elem annak meghatározásában, hogy egy adott felhasználás belefér-e a tisztességes felhasználás kereteibe. Lásd Stim, Rich: What Is Fair Use? Stanford Libraries. Elérhető: <https://fairuse.stanford.edu/overview/fair-use/what-is-fair-use/>

²³⁰ Lemley, A. Mark - Bryan Casey (2020): Fair learning. *Texas Law Review*. 99, 743.o.

²³¹ Az Authors Guild írói szövetség nevében 17 neves író - akik között van többek között George R.R. Martin, a Trónok harca szerzője is - pert indított az OpenAI ellen a szerzői jogok és a szellemi tulajdon megsértése miatt. Az Authors Guide vezérigazgatója úgy fogalmazott: A kiemelkedő könyveket általában azok írják, akik karrierjüket, sőt életüket azzal töltik, hogy tanulják és tökéletesítsék mesterségüket. Irodalmunk megőrzése érdekében a szerzőknek rendelkezniük kell azzal a lehetőséggel, hogy kontrollálhassák, hogy műveiket a generatív mesterséges intelligencia felhasználhassa-e, és ha igen, azt milyen módokon tegye. (saját fordítás) Lásd: Italie, Hillel (2023): John Grisham, George R.R. Martin and other authors sue OpenAI for copyright infringement. Los Angeles Times. Elérhető: <https://www.latimes.com/world-nation/story/2023-09-20/john-grisham-george-r-r-martin-and-other-authors-sue-openai-for-copyright-infringement> (2023. 09. 20.)

használták fel. A felügyelet szerint a cégcsoport megsértette 2 évvel ezelőtt vállalt kötelezettségeit²³², amikor a Bard nevű LLM chatbotját - amelyet azóta Gemini névre kereszteltek át - a kiadók és hírügynökségek tartalmain képezték ki anélkül, hogy erről értesítették volna őket.²³³

A Nagy Nyelvi Modellek tanítására használt adatkészletek kapcsán a képzés során feldolgozott személyes adatok kérdése szintén érzékeny terület. Az olasz adatvédelmi hatóság, a Garante, egy 2023. március 30-i döntésében szólította fel az OpenAI-t, hogy állítsa le az olasz állampolgárok személyes adatainak ChatGPT által történő feldolgozását, ameddig be nem fejeződik az arra irányuló vizsgálat, hogy cég gyakorlatai maradéktalanul megfelelnek-e az EU GDPR előírásainak. A főbb aggályok között szerepelt, hogy a ChatGPT felhasználói nem kaptak tájékoztatást arról, hogy személyes adataikat milyen célokra használják fel, valamint az is, hogy a modellek által generált kimenetek könnyen tartalmazhatnak előzőleg begyűjtött személyes adatokat.²³⁴ E témakörhöz kötődik, hogy a nyelvi modellek esetén a "felejtéshez való jog"²³⁵ biztosítása komoly nehézségekbe ütközhet. Annak ellenére, hogy a modellek által tárolt és feldolgozott hatalmas mennyiségű tanítóadatok között személyes adatok is előfordulhatnak, ezek célzott eltávolítása technikailag bonyolult, mivel a tanítóadatkészletek gyakran összefonódik és elválaszthatatlan a modell teljes működésétől.

A modellek betanításához szükséges korlátozottan elérhető, változó minőségű, vagy jogellenesen begyűjtött tanítóadatok problémájára lehetséges megoldásul szolgálhatnak az ún. szintetikus tanítóadatok. E kifejezés olyan adatokat jelöl, melyek bár mesterségesen lettek létrehozva, valós adatok statisztikai tulajdonságaival bírnak. Ezek használata elméletben lehetővé teheti a hozzáférést megközelítőleg korlátlan mennyiségben rendelkezésre álló és

²³² A versenyhatóság már 2021 júliusában 500 millió eurós bírságot szabott ki a Google-re, azonban ez a vita 2022-ben megoldódni látszott, amikor a cég visszavonta fellebbezését a bírság ellen és több olyan kötelezettséget is vállalt, ami a felek közötti jövőbeli tárgyalások tisztességességének és átláthatóságának biztosítására, valamint a tartalmakért történő megfelelő díjazás fizetésére irányult.

²³³ Satariano, Adam (2024): France Fines Google Amid A.I. Dispute With News Media. The New York Times. Elérhető: https://www.nytimes.com/2024/03/20/business/france-google-fine.htmlutm_source=www.mindstream.news&utm_medium=newsletter&utm_campaign=france-sues-google (2024. 03. 20.)

²³⁴ Curreli, Eleonora (2023): ChatGPT Case: How the Italian Data Protection Authority Is Trying To Address AI Risks. MONDAQ. Elérhető: <https://www.mondaq.com/italy/privacy-protection/1317670/chatgpt-case-how-the-italian-data-protection-authority-is-trying-to-address-ai-risks> (2023. 05. 08.)

²³⁵ A felejtéshez való jog az Európai Bíróság 2014. május 13-án, a "Google Spain SL, Google Inc. v Agencia Española de Protección de Datos, Mario Costeja González" ügyben meghozott ítéletére vezethető vissza, amely kimondta, hogy az internetes keresőmotorok üzemeltetőjének - bizonyos feltételek fennállása esetén - teljesíteniük kell az egyének azon kérését, hogy a nevükön keresési eredményként elérhető, szabadon hozzáférhető weboldalakhoz vezető linkeket eltávolítsák. Ezt a jogot a GDPR 17. cikke is megállapítja (törléshez való jog 17. cikk (1) Az érintett jogosult arra, hogy kérésére az adatkezelő indokolatlan késedelem nélkül törölje a rá vonatkozó személyes adatokat, az adatkezelő pedig köteles arra, hogy az érintettre vonatkozó személyes adatokat indokolatlan késedelem nélkül törölje, ha az alábbi indokok valamelyike fennáll: a) a személyes adatokra már nincs szükség abból a célból, amelyből azokat gyűjtötték vagy más módon kezelték;

megbízható adatkészletekhez, miközben csökkenti azok beszerzésének költségeit, időigényét és jogi kockázatait. Azonban a nem valós és nem élő tanítóadatokon alapuló modellképzés a rövidtávú előnyök ellenére, lényeges sebezhetőséget és hosszútávú kockázatot hordhat magában. Abban az esetben, ha egy modell a saját maga által generált tartalmakon tanul egy idő után elveszítheti a képességét arra, hogy változatos és gazdag tartalmat generáljon, mivel a tanulási folyamat során egyre inkább a saját korábban generált tartalmaihoz igazodik, ami csökkenti a diverzitást és az újító képességet. Ez a folyamatot hívják modell összeomlásnak (modell collapse).

5.5.2 Modell összeomlás

A modell összeomlás – amikor az GenMI modell elveszíti képességét az új, kreatív vagy releváns kimenetek generálására - komolyan vehető fenyegetést jelent a nagy nyelvi modellek hosszú távú fenntarthatóságára nézve.²³⁶ Azon túl, hogy a modell elkezd ismétlődő vagy egyhangú válaszokat adni az is problémát jelent, hogy nem tudja megfelelően kezelni az új bemeneti adatokat, melynek oka, hogy az adott modell - tanító adathalmazok formájában - fokozatosan egyre inkább csak a saját "hangját" hallja vissza, torzítva ezzel a tanulási folyamatot.²³⁷ Az internet - és ezáltal a tanítóadatok - homogenizálódásának veszélye különösen a 2022-es év végétől fokozódik, ahogy az LLM modellek szabadon hozzáférhetővé tételével egyre nagyobb mennyiségben árasztották el szöveges szintetikus tartalmak az online felületeket.²³⁸

²³⁶ Mok, Aaron (2023): A disturbing AI phenomenon could completely upend the internet as we know it. Business Insider. Elérhető: <https://www.businessinsider.com/ai-model-collapse-threatens-to-break-internet-2023-8> (2023. 08. 30.)

²³⁷ E jelenségre – tehát amikor a GenMI modellek által korábban létrehozott szintetikus tartalmakat a tanítási folyamat során visszatáplálják egy adott rendszerbe, a kimenetek minőségének romlását előidézve - MAD-ként (Model Autophagy Disorder) is szokás hivatkozni. Alemohammad és társai egy 2023-as tanulmányukban, leírják, hogy a MAD hatására a generált tartalmak egyre kevésbé változatosak, tartalmuk tekintetében kiszámíthatóbbak, unalmasabbak és rosszabb minőségűek lesznek. A kutatók kísérletei során modell kimenetei öt képzési ciklus alatt használhatatlanná váltak. Lásd: Alemohammad, S., Casco-Rodriguez, J., Luzi, L., Humayun, A., Babaei, H.R., LeJeune, D., Siahkoobi, A., & Baraniuk, R. (2023): Self-Consuming Generative Models Go MAD. arXiv:2307.01850; Dupré H. Maggie (2023): AI Loses Its Mind After Being Trained on AI-Generated Data. Futurism. Elérhető: <https://futurism.com/ai-trained-ai-generated-data> (2023. 12. 07.)

²³⁸ Alig több mint egy év telt el azóta, hogy a fejlett nyelvi modellen nyilvánosan is elérhetővé váltak, de máris előzönlöttek a mesterségesen generált szintetikus tartalmak a webet. Egy 2024-es 404media jelentés szerint a Google News-on fellelhető LLM által generált cikkek száma hónapról hónapra drasztikusan nőtt. Lásd: Cox, Joseph (2024): Google News Is Boosting Garbage AI-Generated Articles. 404media. Elérhető: https://www.404media.co/google-news-is-boosting-garbage-ai-generated-articles/?utm_source=www.therundown.ai&utm_medium=news-letter&utm_campaign=zuck-plans-to-build-open-source-agi (2024. 01. 18.)

Shumailov és társai a GenMI modelleke vizsgálva két összeomlás típust különböztetnek meg, melyeket (1) korai összeomlásnak és (2) előrehaladott összeomlásnak neveznek (early and late model collapse).

(1) A korai modell összeomlás olyan helyzetet jelöl, amikor egy GenMI modell nem képes megőrizni azokat az információkat, amelyek csak ritkán fordulnak elő a tanító adatokban. Ezeket az információkat gyakran a tanítóadatok "farok" részének nevezik, ami az adateloszlás szélsőséges értékeit jelöli (tails of the distribution). Ez a probléma tipikusan a tanulási folyamat korai szakaszában jelentkezik, amikor a modell még csak most kezd saját generált adatain alapuló tanulásra áttérni. Ennek eredményeként a modell által létrehozott adatok kevésbé tükrözik hűen a valóságot, és a ritkább információk fokozatosan eltűnnek a generált tartalomból.²³⁹ Például nyelvi modellt irodalmi szövegek létrehozására képeznek a. A tanítási adatok tartalmazhatnak ritkán előforduló különleges szavakat vagy kifejezéseket. Kezdetben a modell még képes lehet ezeket megtanulni és használni, de ahogy egyre többet támaszkodik saját, előállított szövegeire a tanulás során, előfordulhat, hogy ezek a ritka elemek fokozatosan kezdenek eltűnni a modell által generált szövegekből (ezáltal csökkentve a modell képességét arra, hogy változatos és gazdag nyelvi kimenetet produkáljon).

Az előrehaladott modell összeomlás akkor következik be, amikor egy GenMI modell túl hosszú ideig támaszkodik saját maga által generált adatokra, és emiatt a tanulási folyamat során összekeveri az adateloszlásokat. Ez azt jelenti, hogy a modell a különböző adatcsoportokat, amelyeket korábban megkülönböztetett, egy idő után már nem tudja hatékonyan elkülöníteni, és ezáltal egy egységes, homogén kimenetet kezd el előállítani. Ezt a homogenitást nevezik "konvergens pontnak", ami azt jelzi, hogy a modell már nem képes adekvátnan reagálni vagy adaptálódni új vagy változó adatokhoz, mivel minden bemenetet hasonlóképpen kezel. Például egy képgeneráló modellt, amelyet arra képeznek, hogy különböző állatokat ábrázoló képeket hozzon létre. Kezdetben a modell az eredeti tanító adatok alapján tanul, és valósághű állatképeket generál. Azonban ahogy egyre inkább a saját kimeneteire támaszkodik, előfordulhat, hogy a különböző állatok jellemzői elkezdnek összemosódni, és a modell olyan képeket hoz létre, amelyek nem felelnek meg egyetlen valós állatnak sem.²⁴⁰

Végezetül érdemes szót ejteni egy a nyelvi modellek képzését érintő olyan nehézségről is, ami nem a tanulóadatok, hanem a modellek képzési folyamatának torzításához köthető. A főként

²³⁹ Shumailov, I., Shumaylov, Z., Zhao, Y., Gal, Y., Papernot, N., & Anderson, R. (2023): The Curse of Recursion: Training on Generated Data Makes Models Forget. 3-6.o. ArXiv, abs/2305.17493.

²⁴⁰ Shumailov, I., Shumaylov, Z., Zhao, Y., Gal, Y., Papernot, N., & Anderson, R. (2023): The Curse of Recursion: Training on Generated Data Makes Models Forget. 6-12.o. ArXiv, abs/2305.17493.

hosszabb bemenetek feldolgozásánál és kimenetek generálásnál jelentkező ún. *lost in the middle* (közép elvesztése) problémája, abból ered, hogy a nyelvi modellek hajlamosak jobban figyelembe venni a szövegkontextus elején található információkat, mint a középső vagy végső részeket (pozíciós elfogultság). Ez azért fordul elő, mert a modellek finomhangolása során a kezdeti utasítások nagyobb súllyal esnek latba a tanulási és válaszadási folyamatban, melynek eredményeként a modellek alábecsülik a szöveg közepén vagy végén található releváns adatokat.²⁴¹ Ez különösen akkor okoz problémát, amikor nagy mennyiségű adatot kell feldolgozni, és a lényeges információk egyenlőtlenül oszlanak el a szövegben (például szakmai dokumentumok összefoglalása). Az ilyen típusú elfogultság azáltal, hogy hibás, hallucinációtól terhes, vagy a kontextusra nézve irreleváns tartalmakat hoz létre szintén korlátozza a nyelvmodellek alkalmazhatóságát olyan helyzetekben, ahol a pontosság és az információk teljes körű értékelése kritikus fontosságú.²⁴²

5.6 Szintetikus vagy valódi szöveges tartalom?

Ahogy arra több a közelmúltban közzétett tanulmány is felhívta a figyelmet az emberek egyre kevésbé képesek azonosítani és megkülönböztetni a mesterségesen előállított szöveges tartalmakat.²⁴³ Ennek orvoslására érvényes stratégia lehet MI technológiákat segítségét kérni.²⁴⁴ A különböző MI-detektorok alkalmazása azonban hamar rámutatott arra sebezhetőségre, hogy ezek az eszközök paradox módon maguk is hozzájárulhatnak a mesterséges tartalmakat előállító MI fejlődéséhez. A fejlesztők ugyanis ezeket a detektorokat használják arra, hogy teszteljék és finomítsák azokat az algoritmusokat, amelyeket az MI-detektorokat megkerülő, mesterséges

²⁴¹ Liu, N.F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2023): Lost in the Middle: How Language Models Use Long Contexts. *Transactions of the Association for Computational Linguistics*, 12, 157-173.o.

²⁴² Több próbálkozás is született a helyzet orvoslására, melyek közül jelenleg az (IN2) tűnik a legígéretesebbnek. Az új INformation-INtensive (IN2) adatvezérelt képzési módszer hosszú kontextusból álló kérdés-válasz párokat használ, ahol a válaszok olyan kritikus információkra épülnek, melyek véletlenszerűen elhelyezett rövid szegmensekben találhatók. Ezáltal a modellek képesek kifinomultabban érzékelni az információkat, és integrálni azokat a hosszú kontextus különböző részeiről, ami jelentősen javítja azok teljesítményét hosszú szövegek feldolgozásában (lényegében IN2 képzés megtanítja a modellt, hogy a lényeges információk a hosszú kontextus bármely pontján megtalálhatóak lehetnek, nem csak a szélein). Az új képzési megközelítésről lásd bővebben: Shengnan An, Zexiong Ma, Zeqi Lin, Nanning Zheng, Jian-Guang Lou (2024): Make Your LLM Fully Utilize the Context. 3-10.o. arXiv:2404.16811

²⁴³ N. C. Köbis and L. Mossink (2021): Artificial intelligence versus Maya Angelou: Experimental evidence that people cannot differentiate AI-generated from human-written poetry. *Computers in Human Behavior*, vol.114(2):106553 doi: 10.1016/j.chb.2020.106553

²⁴⁴ Gyakran alkalmazott módszerek például a digitális vízjelek beillesztése, az időbélyegzés(timestamping) a metaadat alapú azonosítás, vagy a blockchain technológia.

tartalmak létrehozására terveznek.²⁴⁵ Ennek eredményeképpen minden új detektálási technológia előrelépése potenciálisan újabb, fejlettebb generatív MI megoldásokat inspirál, amelyek képesek kijátszani a meglévő biztonsági intézkedéseket. Ez az öngerjesztő folyamat — ahol minden új detekciós technológia fejlesztése újabb, kifinomultabb MI megoldás születését ösztönzi — egyfajta fegyverkezési versenyt idéz elő.²⁴⁶ Itt minden technológiai előrelépés nemcsak a biztonságot növeli, de hozzájárul a biztonsági intézkedések megkerülésére képes új technológiák fejlesztéséhez is.

Többen felhívták arra a kockázatra a figyelmet, hogy a szintetikus szöveges tartalmak azonosítása egyre bizonytalanabbá válik.²⁴⁷ Ennek oka az MI- detektorok pontatlanságával hozható összefüggésbe, amit jól szemléltet az is, hogy néhány hónappal ezelőtt az OpenAI azt javasolta felhasználóinak, hogy az általuk 2023-ban elérhetővé tett tartalomkategorizáló szoftver következtetéseit főszabályként csak kiegészítő információként értelmezzék, és ne támaszkodjanak kizárólag ezekre az eredményekre az MI által generált tartalmak azonosításához.²⁴⁸

5.7 (Gen)MI-t hoz a jövő?

²⁴⁵ Lim, N., Kuan, M., Pu, M., Lim, M., & Chong, C. (2022): Metamorphic Testing-based Adversarial Attack to Fool Deepfake Detectors. *2022 26th International Conference on Pattern Recognition (ICPR)*, 2503. <https://doi.org/10.1109/ICPR56361.2022.9956543>.

²⁴⁶ Sugumaran, D., John, Y., C, J., Joshi, K., Manikandan, G., & Jakka, G. (2023): Cyber Defence Based on Artificial Intelligence and Neural Network Model in Cybersecurity. *2023 Eighth International Conference on Science Technology Engineering and Mathematics (ICONSTEM)*, 1-8.o. <https://doi.org/10.1109/ICONSTEM56934.2023.10142590>.

²⁴⁷ Amióta a Nagy Nyelvi Modellek mindenki számára szabadon elérhetővé váltak – kiváltképp az új technológiák iránt nyitott diákok számára – megsokszorozódott mind a technológiát használók, mind pedig a használattal tévesen megvádolt esetek száma. Jelenleg úgy tűnik, hogy az oktatási intézmények által leggyakrabban használt szoftvereknek (DeepTrace / Sensity / Fakespot / Turnitin) egyre kevésbé meghibázhatók a mesterséges tartalmak azonosításában, elsősorban a növekvő falspozitív értékelések száma miatt. A szoftverek hiányosságairól bővebben: Sadasivan, V., Kumar, A., Balasubramanian, S., Wang, W., Feizi, S. (2023): Can AI-Generated Text be Reliably Detected?. *ArXiv*, abs/2303.11156; Subramaniam, R. (2023). Identifying Text Classification Failures in Multilingual AI-Generated Content. *International Journal of Artificial Intelligence & Applications*. Vol.14(5), 57-63.o. <https://doi.org/10.5121/ijaia.2023.14505>

²⁴⁸ Kirchner JH, Ahmad L, Aaronson S, Leike J (2023): OpenAI: New AI classifier for indicating AI-written text. Elérhető: <https://openai.com/blog/new-ai-classifier-for-indicating-ai-written-text> (2023. 05. 20.). 2023 július 20-án frissítették az oldalt, a következő szöveget hozzáillesztve. *A pontosság és megbízhatóság alacsony szintje miatt az LLM szövegazonosító a továbbiakban nem elérhető. A visszajelzéseket feldolgozva jelenleg hatékonyabb tartalommeghatározási technikákat kutatunk a szintetikus szövegek számára. Elköteleztük magunkat amellelt, hogy a jövőben kifejlesztünk és bevezetünk olyan mechanizmusokat, amelyek lehetővé teszik a felhasználók számára, hogy megértsék, hogy egy audio vagy vizuális tartalom MI által generált-e.* (saját fordítás):

A GenMI körüli modern diskurzusnak több új irányvonala is megjelent. Ezek közé tartozik (1) a modellek tanításából eredő növekvő – és hosszútávon nem fenntartható – energiaszükséglet és környezetkárosítás kérdése (2) Az emberközi kapcsolatok, illetve az ember-gép közötti kapcsolatok minőségi átalakulása ²⁴⁹ (3) A GenMI modellek munkaerőpiacra gyakorolt hatása (4) Az új MI generációt képviselő autonóm ágensek (InteraktívMI) megjelenésének hatásai

(1) A Carnegie Mellon Egyetem és a Hugging Face tech. vállalat által készített 2023-as tanulmány, hogy a GenMI képgeneráló modellek felhasználói szintű használata rengeteg energiát emészt fel. Egyetlen képnek az előállítás - amely az OpenAI Dall-e alkalmazásában egy körülbelül 5 másodperces művelet - megközelítőleg egy okostelefon akkumulátorának teljes feltöltéséhez szükséges energiamennyiséget használ fel.²⁵⁰ Mivel a hatalmas adatközpontok, melyek az MI forradalmat táplálják, szintén egyre több energiát igényel, a Wells Fargo bankcsoport előrejelzése szerint, 2030-ig a a GenMI további elterjedésével párhuzamosan akár 20 százalékkal is nőhet az Egyesült Államok elektromos energia iránti igénye²⁵¹ Az energiaéhség problémáját fokozza, hogy ez a technológiai forradalom nagyrészt fosszilis energiahordókra épül. A kirívóan magas energiaszükségleteken túl az MI technológiákkal túlzott használata jelentős környezeti lábnyommal és ökológiai teherrel is jár. Csak érzékeltetés képen a Google 2023-as Környezeti Jelentése szerint a cégcsoport 5,6 milliárd gallon vizet használt fel saját szerverei hűtésére²⁵². Emellett az MI fejlesztés és alkalmazás miatt számottevően emelkedik a szén-dioxid kibocsátás. A Massachusettsi Egyetem egyik kutatócsoportja már 2019-ben azt állította, hogy egyetlen MI modell betanítása több mint 626,000 font szén-dioxid-egyenértékű kibocsátást idéz elő, ami közel ötszöröse az átlagos amerikai autó egész életciklusa során keletkező kibocsátásnak (ideértve az autó gyártását is).²⁵³

²⁴⁹ Harari a gép-ember közötti kapcsolat kialakulásának társadalmi hatásain kívül felhívja a figyelmet az ilyen rendszerek kultúraformáló erejére is. A Nagy Nyelvi Modellek immár képessé váltak arra, hogy az emberi képzetet is gyökeresen átformálják azáltal, hogy egészen új kulturális narratívákat, szimbólumrendszereket, törvényeket és hiedelemvilágokat is teremtsenek számunkra. Meggyőzés erejével átveheti a kultúratermelés feletti irányítást is. Lásd bővebben: Harari, Y. N. (2023): Yuval Noah Harari argues that AI has hacked the operating system of human civilisation. The Economist. <https://www.economist.com/byinvitation/2023/04/28/yuval-noah-harari-argues-that-ai-has-hacked-the-operating-system-of-human-civilisation> (2023. 04. 28.)

²⁵⁰ Luccioni, A. S., Jernite, Y., Strubell, E. (2023): Power hungry processing: Watts driving the cost of ai deployment? 13-14o. arXiv preprint arXiv:2311.16863

²⁵¹ Wells Fargo Bank (2023): Investor Presentation - 2023. 10.o. Elérhető: https://www08.wellsfargomedia.com/assets/pdf/about/investor-relations/presentations/2023/april-investor-presentation.pdf?utm_source=www.mindstream.news&utm_medium=newsletter&utm_campaign=ai-vs-energy-consumption (2023. 08. 12.)

²⁵² Google (2023): Environmental Report. 49-55o

²⁵³ Strubell, Emma, A. Ganesh, A. McCallum (2019): Energy and Policy Considerations for Deep Learning in NLP. arXiv:1906.02243

(2) Az LLM rendszerek meggyőző nyelvi képességeik révén már most meglepően szoros és mély kapcsolatokat alakítanak ki a felhasználókkal. Annak ellenére, hogy e rendszereknek sem valódi tudata sem érzelmei nincsenek, az emberek hajlamosak kötődni hozzájuk pusztán a látszólagos intimitás miatt. Ez könnyen kiszolgáltatottságot eredményezhet, kiváltképp ha arra használják, hogy befolyásolják az emberek döntéseit, magatartását. Ahogy arról még a későbbiekben szó lesz, a modern digitális környezet jelenleg a figyelméért és a kattintásokért folytatott küzdelem terepe. Ez egy szempillantás alatt új, mélyebb szintre léphet, ahol a GenMI a felhasználók gondolataiért, érzelmeiért és meggyőződéséért állnak majd harcba.

(3) A GenMI megjelenése jelentős hatással lehet számos szakmára és munkakörre. Bár az évezred elején még az volt az uralkodó vélekedés²⁵⁴, hogy az MI és a robotizáció elterjedése elsősorban a fizikai munkások megélhetését fogja veszélyeztetni, a 2020-as évek elejére kijelenthető, hogy a GenMI *fehérgallérossá vált*.²⁵⁵ Csak néhány olyan területet és foglalkozást kiemelve, melyeket veszélyeztet - vagy alapjaiban alakíthat át a technológia: az írott tartalomgyártás területén dolgozók, mint az újságírók, szövegírók, bloggerek és forgatókönyvírók szembesülhetnek elsőként komoly kihívásokkal, minthogy a nagy nyelvi modellek már most is képesek gyorsan és nagy mennyiségben létrehozni emberi kreativitási szintet elérő szövegeket. De a vizuális területeken is érezhető a hatás: a grafikai tervezők, illusztrátorok, logótervezők és web-designerek munkája is átalakulhat a képgeneráló GenMI fejlődésével. A nyelvi készségeket igénylő munkakörök sem maradnak érintetlenek, a fordítók és tolmácsok szerepe változhat, ahogy az LLM fordítóprogramok egyre pontosabbá válnak (de ide értendő az irodalmi lektorok és szerkesztők munkája is). Az ügyfélszolgálati szektorban is várható átalakulás, mivel a chatbotok és virtuális asszisztensek képesek lehetnek átvenni bizonyos feladatokat. A jogi területen a jogi asszisztensek és kutatók munkája részben

²⁵⁴ Például Martin Ford is erről ír A robotok kora c. könyvében. Érdekeség, hogy ugyan legtöbbször egyetértettek Ford álláspontjával a robotikai fejlesztések munkaerőpiacra gyakorolt hatásával kapcsolatosan, az ún. Moravec paradox már az 1980-as években viszonylag pontos előrejelzést adott az MI technológiai fejlődésének ívéről. E kifejezés a robotika és az MI kutatás azon meglepő következtetését írja körbe, mely értelmében látszólag egyszerű szenzomotoros tevékenységek (tárgyak megragadása, akadálytalan mozgás, gépi térlátás) végrehajtása lényegesen nehezebb, összetettebb feladat, ráadásul sokkal több számítási erőforrást is igényel annál, mint a rendszereket intelligens képességekkel történő ellátása (absztrakt logikai gondolkodás, magas szintű fordítás, integrálás, sakkjáték szabályainak elsajátítása, tételbizonyítás stb.). Hans Moravec kanadai-amerikai tudós, 1988-ban közzétett könyvében úgy fogalmaz, hogy viszonylag könnyű olyan szintre fejleszteni az MI rendszereket, hogy azok "a felnőttekhez mérhető szintű eredményeket érjenek el intelligenciateszteken vagy a dámapjátékban, de a lehetetlent súrolóan nehéz eljuttatni őket az egy éves gyerekek szintjére az észlelés és manőverező képesség területén". Moravec, Hans (1988): *Mind Children*. Cambridge, Harvard University Press. 15-16.o.

²⁵⁵ A kifejezést Tilesch György MI-kutató használta 2023 szeptember 11-én a Magyar Zene Házában megrendezett az Economx AI Summit konferencián. Lásd: Németh Tamás (2023): Tilesch György: Sokan még mindig a Transformers-filmek szintjén kezelik a mesterséges intelligenciát. EconomX. Elérhető: <https://www.economx.hu/gazdasag/ai-summit-2023-tilesch-gyorgy-mesterseges-intelligencia-kkv-k-munkaltatok.777160.html> (2023. 09. 11.)

automatizálható lehet, különösen az adatfeldolgozás és -elemzés terén. Az adatkezeléssel foglalkozó szakmák, mint az adatelemzők és adatbeviteli munkatársak munkája is átalakulhat a GenMI képességeinek köszönhetően. A programozás területén, főleg az egyszerűbb kódolási feladatok esetében, szintén várható változás. Szintén nagy változás várható a kreatív iparágakban: a digitális marketing területén a közösségi média menedzserek és SEO szakértők munkája nagyrészt kiválthatóvá válik, ahogy az MI-alapú eszközök egyre inkább képesek lesznek tartalmak optimalizálására és elemzésére.

(4) A Google DeepMind társalapítója, Mustafa Suleyman úgy véli, hogy a GenMI csak rövid, átmeneti szakaszt jelent a technológiaicsalád fejlődésében. A következő hullám az autonóm és célirányos cselekvésekre képes interaktív MI lesz. Amíg a jelenlegi generatív modellek bemeneti adatokból új adatokat előállítására képesek az interaktív MI ezen túl azt is lehetővé teszi majd az emberek számára, hogy különböző feladatokat is rábízzanak²⁵⁶Egy személyre szabott digitális asszisztenshez hasonlóan képes lesz a felhasználók eszközein másokkal interakcióba lépni, és különböző feladatokat végrehajtani (útiterv-készítése, repülőjegy-foglalás, élte-rendelés stb.) Ezek mindegyike számtalan új szabályozási kérdést vet fel.

Az elmúlt másfél évben azonban konszenzus látszik kialakulni abban, hogy a GenMI modellek legégetőbb, rövid távú kihívását az jelenti, hogy a korlátlan mennyiségű, olcsón előállítható szintetikus szöveges és audiovizuális tartalmat a dezinformációs- és spamkampányoktól kezdve, a pénzügyi csaláson²⁵⁷ és személyazonosság-lopáson keresztül, a politikai manipuláció és választási befolyásolásig bezárólag rengeteg mindenre fel lehet használni.²⁵⁸ Azonban a

²⁵⁶ Az interaktív MI/autonóm ágensek olyan rendszerek, amelyek képesek önállóan döntéseket hozni és cselekedni anélkül, hogy közvetlen emberi beavatkozásra lenne szükségük. Bővebben: Suleyman, Mustafa, Michael Bhaskar (2023): A következő hullám. Mesterséges intelligencia, technológia, hatalom és a 21. század legnagyobb kihívása. (ford. Farkas Veronika). Budapest, Magnólia kiadó. Will Douglas Heaven (2023): DeepMind's cofounder: Generative AI is just a phase. What's next is interactive AI. MIT Technology Review. Elérhető: <https://www.technologyreview.com/2023/09/15/1079624/deepmind-inflection-generative-ai-whats-next-mustafa-suleyman/> (2023. 09. 15.)

²⁵⁷ Annak egyik lehetséges következményét, hogy gyakorlatilag mindenki számára lehetségessé vált, a magas minőségű szöveges tartalmak nagy mennyiségben történő előállítása jól érzékelteti, hogy a ChatGPT 2022 novemberi megjelenése óta több, mint 1.260%-kal nőtt a hitelesnek tűnő, tehát kevesebb nyelvtani hibát tartalmazó és koherens szövegszerkezettel bíró (ezáltal hatékonyabb és eredményesebb) adathalász e-mailek száma. Lásd: Nelson, Jason (2023): Email Phishing Attacks Up 1,265% Since ChatGPT Launched. EMERGE. Elérhető: <https://decrypt.co/203564/since-chatgpt-launch-phishing-emails-are-up-1265-slashnext>

²⁵⁸ 2023 januárjában számolt be arról a New York Times, hogy az OpenAI felfüggesztette egy olyan felhasználójának fiokját, aki a demokrata elnökjelölt-aspiráns Dean Phillips képviselőt megszemélyesítő automatizált botba integrálta a cég nagy nyelvi modelljét. A *Dean Bot*, amely valós időben tudott kommunikálni a szavazókkal egy weboldalon keresztül, ugyan tartalmazott figyelmeztetést arra vonatkozóan, hogy nem a valódi Dean Phillips, az OpenAI mégsem engedte annak alkalmazását. A cég azzal indokolta a döntést, hogy belső szabályzatuk tiltja a technológiájuk politikai kampányokban történő olyan jellegű felhasználását, amely során az érintett beleegyezése nélkül visszaélnek annak a személyazonosságával. Dwoskin, Elizabeth (2024): OpenAI suspends bot developer for presidential hopeful Dean Phillips. The Washington Post. Elérhető:

GenMI modellek által létrehozott szintetikus tartalmak példátlan gyorsaságú elterjedése hosszútávú hatással is bír majd a körülöttünk lévő online és offline információs közegre. Ennek feltérképezésére vállalkozik a következő fejezet. A folyamat megértéséhez egészen az internet létrejöttének kezdetéig érdemes visszamenni.

https://www.washingtonpost.com/technology/2024/01/20/openai-dean-phillips-ban-chatgpt/?utm_source=www.therundown.ai&utm_medium=newsletter&utm_campaign=sam-altman-s-secret-ai-chip-venture (2024. 01. 20.)

6 Szép új (online) világ

Orwell vagy Athén?²⁵⁹

Az internet elterjedését kezdetben általános optimizmus övezte. E pozitív hangulat abból a reményből fakadt, hogy a világháló 1990-es évektől kezdődő²⁶⁰ térnyerésével megindulhat majd az információ és tudás egyetemes demokratizálódása. A hagyományos információáramlási struktúrák átalakulásával a történelem során felhalmozott ismeretek széles körben elérhetővé és hozzáférhetővé váltak. A világháló új csatornákat biztosított az információ megosztására lebontva azon időbeli, térbeli és anyagi korlátokat, amelyek korábban akadályozták a tudáshoz való egyetemes hozzáférést. Az országok, kormányok, állampolgárok közötti kommunikációs korlátok csökkenésével az emberek a világ legtávolabb pontjain is könnyen kapcsolatba léphettek egymással, elősegítve ezzel a kulturális információs cserét, és a nemzetközi együttműködést.²⁶¹ Egy rövid ideig úgy tűnt, hogy az információ demokratizálódásán túl az internet aktív szerepet játszhat a politikai demokratizálódási folyamatokban is. A különböző online platformok és a közösségi média - mint felszabadító technológia (liberation technology)²⁶² - lehetőséget teremtettek az állampolgároknak arra, hogy éljenek információs szabadságukkal, önként szerveződjenek, és szabadon fejezzék ki véleményüket ott, ahol a hagyományos média korlátozott vagy szoros állami ellenőrzés alatt állt. Ennek jelentősége különösen az arab tavasz eseményei során vált nyilvánvalóvá, amikor a Facebook és a Twitter központi szerepet játszott a legtöbbször alulról szerveződő, reformpárti tüntetések és politikai mozgalmak szervezésében és koordinálásában²⁶³

²⁵⁹ van de Donk, W. B. H. J., & Tops, P. W. (1995): Orwell or Athens? Informatization and the Future of Democracy. In W. B. H. J. van Donk, I. T. M. Snellen, & P. W. Tops (szerk.): Orwell in Athens: a Perspective on Informatization and Democracy. IOS Press. 13-32.o.

²⁶⁰ Tim Berners-Lee 1989-ben hozta létre a máig leghíresebb „World Wide Web” globális információs rendszert, amely a különböző dokumentumokat ún. hiperhivatkozással kapcsolta össze. Ugyan a rendszert 1991-ben már használták kutatók, a szélesebb közönség számára csak a 90-es évek első felében, a webböngészők elterjedésével vált ismertté. (ezek közül a leghíresebb az 1993 április 22-én megjelent Mosaic böngésző).

²⁶¹ Az internet kezdeti fejlődéséről lásd bővebben: Castells, Manuel. (2002): Az Internet-galaxis (Gondolatok internetről, üzletről és társadalomról). Budapest, Network Twenty One Hungary kiadó.

²⁶² A liberation technology kifejezés népszerűsítése Larry Diamond amerikai politikai szociológushoz köthető. Diamond a kifejezést annak leírására alkalmazta, hogy miként tudják az információs és kommunikációs technológiák (mint például a közösségi média) elősegíteni a demokratikus folyamatokat, az emberi jogok védelmét és a politikai aktivizmust az autoritás és elnyomás elleni harcban. Bővebben lásd: Diamond, Larry (2010): Liberation Technology. *Journal of Democracy*. vol. 21, no. 3, 69-83.o.

²⁶³ A fentebb említett online platformok és közösségi médiumok (Facebook, X) nagyban segítették a tüntetéshullámban érintett országokban a kollektív fellépés sikeres megszervezését. Ez köszönhető volt többek

Azonban az online tér iránt érzett kezdeti lelkesedés hamar szertefoszlott. Az idő előrehaladtával az internet egyre kevésbé testesítette meg a születésekor vizionált szabad információáramlást biztosító, korlátozások és beavatkozásoktól mentes demokratikus környezetet. Napjainkban az online térben történő információgyűjtést, a digitális tartalmakhoz való hozzáférést kormányok vagy vállalatok által működtetett algoritmus alapú döntések szabályozzák, amelyek irányát - kevés átláthatóság és közösségi felügyelet mellett - gyakran a hatalom kiteljesítése vagy a profit maximalizálása szabja meg. Az információs közeg ahelyett, hogy tájékozottabbá tette volna a felhasználókat sok esetben kételyt ébreszt és dezinformál, ahelyett, hogy közelebb hozta volna az embereket inkább ellentéteket gerjeszt és polarizál. Ahhoz, hogy ennek az okait és következményeit megértsük, az online környezet fejlődését és jelenlegi működését kell górcső alá venni.

6.1 Zajos terek, kétes információk

Az internet megjelenésének kezdetén, - a Web 1.0 korszakában - a weboldalak statikus jellegűek voltak, főként egyszerű HTML forráskódból álltak, így a felhasználók számára kizárólag előre megírt szöveges és képi tartalmakat jeleníthettek meg. E webhelyek csupán egyirányú kommunikációra nyújtottak lehetőséget: a tartalmat a weboldal üzemeltetői helyezték el, míg az oda látogatók pusztán fogyasztói lehettek a közreadott információknak, aktív interakciós vagy hozzászólási lehetőség nélkül. A felhasználó tehát hozzáférhetett az információhoz, de nem oszthatta meg saját bemenetét. A ma ismert keresőmotorok még gyerekcipőben voltak, a közösségi média pedig nem is létezett. A kétezres évek elején jelent meg az internet új generációja - a web 2.0 - ami radikálisan átformálta a digitális világot. A Web 2.0 már egy dinamikus, felhasználóközpontú digitális környezetté alakult, ahol a tartalom már nem csak olvasható, hanem szerkeszthető, bővíthető és módosítható is egyben. A kommunikációs csatornák kétirányúvá váltak a blogokon, fórumokon és komment szekciókon keresztül.

között annak, hogy gyorsan és valós időben voltak képesek tömegek számára információt terjeszteni a szervezett tüntetésekről, a hatóságok közbeavatkozásáról, a várható rendőri jelenlétről stb. Segítségükkel széles közönséghez juthattak el az adott ügyekkel kapcsolatos hírek, és olyan embereket is mozgósíthattak, akiket hagyományos eszközökkel talán nem értek volna el (a helyi és regionális résztvevőkön túl, a nemzetközi sajtót, megfigyelőket és egyéb támogatókra is ideértve). A témáról lásd bővebben: Howard, Philip N. és Muzammil M. Hussain (2013): *Democracy's Fourth Wave? Digital Media and the Arab Spring*. Oxford University Press; Tufekci, Zeynep. (2017): *Twitter and Tear Gas: The Power and Fragility of Networked Protest*. Yale University Press

Megjelentek a könnyen kezelhető tartalomkezelő rendszerek (CMS), amelyek lehetővé tették a felhasználók számára a tartalom egyszerű közzétételét és szerkesztését, az egyszerű szövegalapú HTML weboldalak mellett elterjedtek a multimédiás tartalmak, (pl. audiók, GIF-ek, videók), a keresőmotorok pedig hatékony eszközzé váltak a webes tér navigálásához.²⁶⁴ A közösségi média megjelenése forradalmasította az emberek közötti interakciók formáit, intenzitását és gyakoriságát. Ez az új korszak egy sokkal interaktívabb környezetet teremtett, melyben a blogok, közösségi hálózatok, és videómegosztó oldalak, lehetővé tették, hogy bárki saját tartalmat hozzon létre és osszon meg.²⁶⁵ Ennek eredményeképpen összeesődtak a határok a tartalmak fogyasztása és előállítása között, a felhasználó egyidejűleg tartalomfogyasztóvá és tartalomgyártóvá is vált.²⁶⁶ A web 2.0 elterjedésével a hagyományos kapuőri²⁶⁷ szerep is gyökeres változáson ment keresztül mely okán egy rövid ideig szinte bárki közzé tehetett az általa létrehozott tartalmat anélkül, hogy ezt hagyományos szerkesztői folyamatok szűrték volna.²⁶⁸ Kétségtelen, hogy ez az új online környezet jelentősen bővítette az elérhető információk körét, de egyben komoly kihívást is állított a digitális tartalmak hitelességének és a minőségének megőrzése kapcsán.

A kétezres évek elején a korlátozásoktól mentes - de számottevő információs zajtól terhes - digitális környezetet fokozatosan alakították át azon online közvetítők megjelenése, amelyek a tartalomszolgáltatók és a felhasználók közötti kapcsolatként egyre jelentősebb szerepet játszottak

²⁶⁴ O'Reilly, T. (2005): What Is Web 2.0: Design Patterns and Business Models for the Next Generation of Software. *Communications & Strategies*, No.65 (1), 17-37.o

²⁶⁵ Szőke Gergely László (2012): Az Európai adatvédelmi jog megújítása. Tendenciák és lehetőségek az önszabályozás területén. PTE-ÁJK, 57-58.o.

²⁶⁶ E jelenség leírására gyakran alkalmazzák az Alvin Toffler amerikai író és jövőkutató által 1980-ban bevezetett prosumer kifejezést, amely a "producer" (gyártó) és "consumer" (fogyasztó) szavak összevonásából származik, és a fogyasztói szerepkör megváltozását jelenti, ahol a fogyasztók aktívan részt vesznek a termékek és szolgáltatások létrehozásában. Lásd: Toffler, Alvin (1980): *The Third Wave*. New York, Bantam Books. A későbbiekben még utalás szintjén szó lesz arról is, miként segíti a felhasználók kettős szerepe a dezinformáció, az álhírek, az áltudományos teóriák és konspirációs elméletek terjedését.

²⁶⁷ A kapuőr (gatekeeper") fogalma a kommunikációtudomány területén olyan személyt, jogi személyt takar, aki vagy amely dönt arról, hogy mely információk jutnak el a szélesebb közönséghez. A kapuőri szerep a hagyományos médiában –nyomtatott sajtóban, televízióban, rádióban – még leggyakrabban az újságírókhoz, kiadványszerkesztőkhöz, vagy hírigazgatókhoz kötődött. Ők döntöttek arról, hogy mely hírek, cikkek vagy jelentések érdemelnek figyelmet, és hogyan kerülnek azok bemutatásra. Ez a szerep az internet 2.0 kezdetekor egy rövid ideig demokratizálódni látszott, majd irányítását átvették a piacvezető tech. vállalatok által működtetett közvetítő szolgáltatások (pl: keresőmotorok, közösségi média platformok)

²⁶⁸ Ezzel kapcsolatban Török Bernát így fogalmaz; „*A közösségi média – működésének első időszakában – megteremtette a társadalmi párbeszéd kapuőrök nélküli szféráját, ahol nem lap- vagy tévétulajdonosokon, nem is szerkesztőségeken, de még csak nem is újságírókon, hanem mindenekelőtt a megszólalón múlik, hozzászól-e közügyeinkhez.*” Lásd: Török Bernát (2022): A szólásszabadság a közösségi platformokon és a Digital Services Act. In: Török Bernát és Zódi Zsolt (szerk.): *Az internetes platformok kora*. Budapest, Ludovika Egyetemi Kiadó, 197.o.

annak eldöntésében, hogy kihez, miként és milyen információ jutott el az online térben (illetve, hogy melyek maradjanak rejtve vagy váljanak nehezen hozzáférhetővé).

Ezek a közvetítők, mint például a közösségi média platformok és a keresőmotorok, jellemzően algoritmusok segítségével szűrték, rangsorolták és szabják személyre a tartalmakat. Az ilyen ajánlóalgoritmusok döntéseikben nagyban támaszkodnak az online személyiségprofilokra, amelyeket a fokozatosan *digitalizálódó egyén* internetes aktivitásából származó adatokból építenek fel.²⁶⁹ Az, hogy milyen módon működik az algoritmus, amely dönt a felhasználó számára hozzáférhetővé tett tartalmakról, kizárólag e rendszerek üzemeltetőin múlott.²⁷⁰ A felhasználók ebből eredő kiszolgáltatottságát pedig a piaci környezetre jellemző koncentrálódás csak tovább növeli.²⁷¹

Ez az új online ökoszisztéma egy olajozottan működő digitális gazdasági modellé fejlődött. Ebben a környezetben a felhasználók információs és tájékozódási szabadságának korlátozása nem cél, csak következmény. Elsődleges célként a felhasználók - vélt, valós, olykor tudatosan

²⁶⁹ . E felhasználói profilok ahogy egyre növekvő adatforrásból táplálkoznak úgy válnak lényegesen összetettebbé is. Lásd bővebben: Pataki Gábor – Szőke Gergely László (2017): Az online személyiségprofilok jelentősége – régi és új kihívások. *Infokommunikáció és jog* 2017/2., 63–70.o.

²⁷⁰ Bár az algoritmikus ajánlórendszerek különböző adatokat és számolási modelleket használhatnak, ezek közül három terjedt el leginkább a gyakorlatban: (1) Tartalomszűrésen alapuló (content-based filtering) módszer, amely a felhasználó korábbi preferenciáiból és az ajánlásra szánt tételek tulajdonságaiból kiindulva tesz javaslatokat. Például egy zene-streaming szolgáltatató platform, ahol a felhasználó korábban főleg klasszikus zenét hallgatott, erre alapozva fog hasonló stílusú, de a felhasználó számára még ismeretlen klasszikus zenei albumokat vagy előadókat ajánlani. Ez a szűrőmodell az tartalmat (mint szöveg, képek vagy hanganyag) elemezve keresi meg a hasonlóságokat, hogy olyan tételeket ajánljon, amelyek illeszkednek a felhasználó által fogyasztott tartalmakhoz. (2) A demográfiai szűrés (demographic filtering) alapja az a feltételezés, hogy a hasonló személyes jellemzőkkel (mint nem, életkor, lakóhely) rendelkező felhasználóknak hasonlóak lehetnek az ízléseik és preferenciáik is. Ennek megfelelően a rendszer a felhasználók demográfiai hasonlóságai alapján tesz javaslatokat. Így például egy fiatal városi nő más ajánlásokat kap, mint egy idősebb vidéki férfi. (3) A kollaboratív szűrés (collaborative filtering) egy adott csoportnál megfigyelt preferenciákra támaszkodik az ajánlások generálásakor. A rendszer azonosítja azokat a felhasználókat, akik hasonló ízlésűek és preferenciákkal rendelkeznek, mint a "aktív" felhasználó, majd az ő értékeléseik és akcióik alapján tesz javaslatokat az aktív felhasználónak. Például, ha egy felhasználó 5 csillagosra értékelte Barcsi Tamás, című könyvét a bookline felületén, akkor a rendszer az adatfeldolgozás során azon vásárlókra összpontosít, akik hasonló értékelést adtak erre a könyvre, és az ő további olvasmányélményeik alapján fog újabb könyveket ajánlani a felhasználónak. Bobadilla, J., Ortega, F. Hernando, A. and Gutierrez, A. (2013): Recommender Systems Survey. Knowledge-Based System. Vol(46), 112-113.o.

²⁷¹ Ezt többek között a digitális gazdaság kapcsán gyakran emlegetett *hálózati hatás* jelensége magyarázza: A ma ismert tech. vállalatok jelentős része azáltal tudta idejében leuralni az online piacot, hogy ők rendelkeztek elsőként számottevő felhasználói bázissal. A domináns helyzetük abból fakadt, hogy a felhasználók nem kerestek alternatív szolgáltatásokat, hiszen a már kiválasztott platformjuk vonzerőjét éppen az adta, hogy az ismerőseik közül sokan szintén azt használták. Ezt erősíti az ún. hólabda hatás is mely szerint, ha egy szolgáltatásnak elegendő felhasználója van sokkal könnyebben válik számára terméke minőségének javítása, ezzel létrehozva egy öngerjesztő folyamatot, amelyben a javuló minőség még több felhasználót vonz, ami további fejlesztést tesz lehetővé és így tovább. Az Uber esetében például minél több sofőr és utas csatlakozik a platformhoz, annál gyorsabb és hatékonyabb a szolgáltatás (rövidebb várakozási idő/több elérhető autó), az eBay online aukciós oldal esetében minél több eladó és vevő használja az oldalt, annál nagyobb a választék és a vásárlási lehetőségek.

formált – preferenciáinak kiszolgálása, és az ebből származó online aktivitás monetizálása szolgál.

6.2 A figyelmed eladó

A figyelemgazdaság korában – ahogy ötletes elnevezése is sugallja – a szolgáltatók közötti verseny a felhasználók figyelmének megragadásáért és megtartásáért zajlik. Ebben a környezetben – ahol a figyelem értékes árucikknek számít – a tartalomszűrő algoritmusok azon tartalmakat rangsorolják előre és teszik láthatóvá, amellyel leginkább képesek lekötni felhasználót. Itt az „ingyenes” szolgáltatásokért a felhasználó figyelmével, interakcióival és az azokból kinyerhető adatokkal fizet. Minél hosszabb idejű és változatosabb a felhasználói aktivitás annál nagyobb a potenciálisan realizálható profit.²⁷² Ez a működési logika szorosan összefügg az olyan népszerű közösségi médiaplatformok kialakításával is, mint a Facebook, a YouTube vagy a TikTok. Ezek a platformok szándékosan aknázzák ki az emberi agy dopamin-választ²⁷³ annak érdekében, hogy fokozzák a felhasználói érdeklődést és interakciókat.²⁷⁴ Amikor a felhasználó pozitív visszajelzést kap a közösségi médiában – legyen az egy kedvelés bejegyzése alatt vagy egy új, automatikusan feldobott videó a TikTokon –, az agyban dopamin szabadul fel, ami kellemes érzést vált ki. Ez a folyamat egyfajta jutalomként működik, ami arra motiválja a felhasználót, hogy folytassa azon tevékenységeket, amelyek ezt a pozitív érzést előidéznek. Ezzel kapcsolatban három pszichológiai hatás kiemelése indokolt (1): A felhasználók egyre több időt töltenek az ilyen közösségi média felületeken, folyamatosan keresve a jutalmazó interakciókat, ami könnyen függőséghez, depresszióhoz és szorongás kialakulásához vezethet.²⁷⁵ (2): A dinamikus és ingergazdag online környezetben a folyamatosan felkínált új tartalmak – mint dopamin stimulánsok – hozzájárulnak ahhoz, hogy a digitális bennszülöttek figyelemküszöbe drasztikusan csökkent. Bár ez kedvez a rövid, egyszerű tartalmak

²⁷² Tim Wu (2016): *The Attention Merchants: The Epic Scramble to Get Inside Our Heads*. New York, Knopf.

²⁷³ A dopamin egy a központi idegrendszerben termelődő neurotranszmitter, amely kiemelt szerepet játszik az agy jutalmazási és motivációs rendszereiben. Rendelkezésre állása (vagy nem állása) nagyban befolyásolja az örömeztet, az elkötelezettséget, és a motivációt. Mivel dopamin rendszer aktiválása jutalmazó érzetet, örömet vagy elégedettséget okoz, hozzájárulhat bizonyos viselkedések megerősítéséhez és ismétléséhez.

²⁷⁴ Ilyen technikák például az értesítések gyakori megjelenítése, a végtelen görgetés, vagy az automatikusan lejátszódó videók.

²⁷⁵ Jonathan Haidt, amerikai pszichológus, a New York University professzora szerint a közösségi média jelentős szerepet játszik a tinédzserek mentális egészségének romlásában. Kutatásai azt mutatják, hogy az elmúlt évtizedben jelentősen megnőtt a depresszió, szorongás, önsértés és az öngyilkossági kísérletek aránya a serdülők körében. Ezek szintje kimutathatóan magasabb azok körében, akik naponta több mint két órát töltenek a közösségi média felületeken. Haidt, Jonathan, Nick Allen (2020): *Scrutinizing the effects of digital technology on mental health*. *Nature*. Vol 578, 226-227o.;

népszerűségének (nem véletlen, hogy a legfiatalabb generáció körében már a TikTok a leglátogatottabb platform), egyúttal korlátozza a hosszabb, összetett, komplex tartalmak megértésének és értelmezésének képességét²⁷⁶ (3): Az új tartalmak folyamatos, automatikus felkínálása egy virtuális "nyúlüregbe" (rabbit hole) taszítja a felhasználókat, amelyben egyre mélyebbre süllyednek egy olyan információs örvényben, ahol gyakran egyre szélsőséesebb tartalmakat fogyasztanak - ezt a jelenséget algoritmikus radikalizációnak is nevezik.²⁷⁷

6.3 A jövőd is eladó?

Shoshana Zuboff amerikai szociálpszichológus 2019-ben megjelent könyvében egy ehhez hasonló gazdasági modellről, a megfigyelési kapitalizmusról ír részletesen. Zuboff arra a következtetésre jut, hogy az a figyelemgazdaság korára még érvényes mondás, miszerint, ha valamilyen szolgáltatás ingyenesen veszel igénybe, akkor te magad vagy a termék, már nem

²⁷⁶ Az információs túlterhelés (information overload) valamint az elmélyedést és alaposabb megértést akadályozó tartalomfogyasztási szokások (kiegészítve a csökkenő figyelemfenntartási képességekkel) könnyen az ún. Dunning-Kruger hatás erősödéséhez vezethetnek. Ez arra a pszichológiai jelenségre utal, miszerint egy adott témában alacsonyabb szintű tudással, gyakorlattal rendelkező személyek rendszerint túlértékelik saját képességeiket vagy esetleges teljesítményüket a területen. A szerzőpáros által lebonyolított kísérlet szerint ez arra vezethető vissza, hogy a személyek a témához kötődő koherens tudásuk hiányában képtelenek voltak olyan önreflexióra, mellyel felismerhetnék az adott témához kapcsolódó tudásbeli, teljesítménybeli korlátokat. A dolgozat témájához kötődve fontos kiemelni azt is, hogy a tudásdeficit nem csak saját magukkal szembeni pozitív elfogultságot okozott, hanem egyúttal azzal is járt, hogy képtelenek voltak felismerni és megkülönböztetni azt, ha egy a témában valóban jártas szakértő gondolatait hallották. Bővebben a kísérletről: Kruger, Justin & Dunning, David. (1999): Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121–1134.o.

²⁷⁷ A "nyúl-üreg hatás" (rabbit-hole effect) olyan automatizált információ-torzulásra utal, mely során a felhasználó egy adott témával kapcsolatban kezd el információkat keresni és tartalmat fogyasztani online, de fokozatosan, lépésről-lépésre a kezdeti témától egyre távolabb eső - és ahhoz egyre kevésbé kapcsolódó - tartalmakhoz jut el az algoritmusok ajánlásai révén. Ez gyakran észrevétlenül történik, és mire a felhasználó feleszmél, már teljesen más témájú tartalmakat fogyaszt, mint amire eredetileg kíváncsi volt. Ez egyfajta akaratlan eltávolodást jelent a kezdeti érdeklődési körtől azáltal, hogy az algoritmusok révén folyamatosan felkínált, de apránként eltérő ajánlások elcsábítják a felhasználót új irányokba. Ezzel összefüggésben megemlítendő, az algoritmikus radikalizáció is, amely azt a jelenséget írja le, hogy az online platformok ajánló algoritmusai hajlamosak fokozatosan egyre szélsőséesebb tartalmakat javasolni a felhasználóknak. Itt megfigyelhető egy trend, ami főként a fiatal férfi felhasználók befolyásolására irányul. Ezt gyakran ártalmatlannak ható mentális és fizikai egészséggel más estekben online pénzszerezéssel és gazdasági tudatossággal kapcsolatos tanácsok közvetítésével valósítják meg a magas követői számmal és eléréssel rendelkező influenszerek. Ebben a nyúlüregben a közösségi médiában egyre nagyobb teret kapó videókat - a fiatal identitáskeresés nehézségeinek monetizálásával főként életmódtanácsadásra összeponstosító tartalmakat - hamar részben vagy egészben felváltják a nemek közötti egyenlőség eszményét, és a szexuális kisebbségek jogait támadó, olykor szélsőjobboldali, sovíniszta vagy antiszemita ideológiákat hirdető videók, melyek a még későbbi szakaszokban az egészen szélsőséges világnézetek felé is terelhetik a felhasználót. A témával kapcsolatos új trendeket – főként a tik-tok videó és a youtube short tartalmakra fókuszálva - lásd bővebben a Dublini egyetem 2024-es tanulmányát: Baker, Catherine, Debbie Ging, Maja Brandt Andreassen (2024): Recommending Toxicity: The role of algorithmic recommender functions on YouTube Shorts and TikTok in promoting male supremacist influencers. DCU Centre, Dublin City University. 2-33o. Elérhető: <https://antibullyingcentre.ie/wp-content/uploads/2024/04/DCU-Toxicity-Full-Report.pdf>

fedile a teljes valóságot. A tényleges helyzet ennél sötétebb képet mutat, hiszen a felhasználók ebből az eleve kiszolgáltatott termék pozícióból mára leértékelődtek egyszerű nyersanyaggá, melynek elsődleges szerepe, hogy a belőle kinyert adatok bemenetként szolgáljanak a jövőbeli magatartások előrejelzésére programozott algoritmusoknak. Ezzel párhuzamosan a valódi termékekké a felhasználók jövőbeli magatartásának előrejelzései váltak. Ezeket az előrejelzési termékeket (prediction production) különböző piaci szereplőknek értékesítik tovább, akik ezáltal hatékonyabban tudják befolyásolni a fogyasztót. Ha a cégek elegendő viselkedési adattal rendelkeznek lehetővé válik az olyan online magatartások ösztönzése is, mint egy adott termék hirdetésre kattintás, videó megnyitása, hírfolyam követése stb. Zuboff szerint bár a megfigyelési kapitalizmus alig pár éve kezdte bontogatni szárnyait, már most is rengeteg ember döntéseire bír komoly befolyással – sok esetben az ő tudtuk nélkül. Ahogy egyre több okosinfrastruktúra kapcsolódik a digitalizált ökoszisztémához, az értékes adathalmazok, az adatelemzés által kínált lehetőségek, valamint a fogyasztói magatartások előrejelzésének és befolyásolásának eszköztárai is növekedni fognak a jövőben.²⁷⁸ Eljön az idő, mikor már lehetetlennek tűnik elképzelni azt, hogy bárki kitörhet ebből a hálóból.

Ebben az új online környezetben a felhasználók nemcsak a piaci szereplők manipulatív gyakorlataival szemben váltak kiszolgáltatottá, hanem a politikai indíttatású befolyásolási technikákkal szemben is. Ennek kifejtése előtt azonban szükséges egy az algoritmikus tartalomajánláshoz kapcsolódó másik jelenségre is kitérni, mely fontos tényezőként játszott szerepet az információs környezet (és egyéni információfeldolgozás) napjainkban tapasztalható eltorzulásában.

6.4 Információs buborékok, személyre szabott valóságok

Eli Parizer 2011-ben megjelent könyvében ismertette meg a médiatudományokban kevésbé jártas szélesebb közönséggel is a szűrőbuborék (filter bubble) fogalmát. A kifejezés arra a jelenségre utal, amikor az online platformokon alkalmazott ajánlóalgoritmusok az internetes aktivitást és hírfogyasztási szokásokat figyelembe véve főként személyre szabott tartalmakat jelenít meg, a felhasználóknak, ezzel egyfajta információs „buborékba” zárva őket. Minél inkább az egyéni preferenciák alapján szűrik és értékelik az információt, annál inkább

²⁷⁸ Diósi Szabolcs, Barcsi Tamás (2021): The legacy of disciplinary society – how relevant is Foucault’s theory today? Évkönyv - Újvidéki egyetem magyar tannyelvű tanítóképző. XVI. évfolyam, 1. szám, 10-33.o.

valószínű, hogy a felhasználók olyan tartalmakhoz férnek hozzá főként, amelyek illeszkednek a meglévő attitűdjeikhez, ismereteikhez és álláspontjukhoz, ezáltal olyan *hiedelemkamrákba*²⁷⁹ szorulnak, melyek megnehezítik számukra, hogy a saját nézeteiktől eltérő információkhoz jussanak.²⁸⁰ Az ilyen szűrők azon túl, hogy meghatározott irányban formálhatják a felhasználó világképét, hozzájárulhatnak az ún. *önbeteljesítő identitások* (self-fulfilling identities) kialakulásához is.²⁸¹ Mivel a felhasználók folyamatosan olyan tartalmakat kapnak, amelyek korábbi érdeklődésüket tükrözik, hamar olyan helyzet állhat elő, ahol az adott egyén múltbeli preferenciái lesznek kihatással a jövőbeli döntéseik, érdeklődési körük, és identitásuk formálódására. Ez az önbeteljesítő hatás korlátozza a felhasználók autonómiáját, mivel kevésbé képesek szabadon választani és felfedezni új vagy kihívást jelentő ötleteket és nézeteket. Ezáltal a szűrőbuborékok nem csak az információáramlást korlátozzák, hanem a sokszínű személyiség fejlődését kialakítását is.

További kihívást jelent, hogy az online környezetben a felhasználók gyakran nem tudatosan választják a szűrőbuborékokat; az algoritmusok automatikusan formálják az információáramlást. Ez új jelenséggé tűnik fel a hagyományos médiafogyasztási szokásokhoz képest, ahol az egyének tudatosan választhatnak olyan újságokat vagy tévécsatornákat, amelyek egy adott ideológiai irányvonalat képviselnek. Ezzel ellentétben az online térben könnyen előfordulhat, hogy a felhasználó nincs is tudatában annak, hogy algoritmusok által szűrt információs környezetben tájékozódik²⁸², amiből egyenesen következik, hogy nem is a saját döntéséből lépett be az információs buborékba (arról nem is beszélve, hogy az abból való kilépés is egyre nehezebbé válik).

6.4.1 Az Információfeldolgozás torzulása

A tartalom alapú szűrési algoritmusok társadalmi hatásait vizsgálva fontos kiemelni a Cass Sunstein által már egy évtizeddel korábban feltárt visszhangkamrákhoz (echo chambers)

²⁷⁹ Veszelszki Ágnes (2021): deepFAKEnews: Az információmanipuláció új módszerei. In Balász László (szerk.): Digitális Kommunikáció és Tudatosság. Budapest, Hungarovox kiadó, 96.o.

²⁸⁰ Veszélyeztetve ezzel az Európai Unió Alapjogi Chartájának 11. cikke, 1. bekezdésében is foglalt információk és eszmék megismerésének szabadságát.

²⁸¹ Eli Pariser (2021): The Filter Bubble. What the Internet is Hiding from You New York. The Penguin Press. 112.-113.o.

²⁸² Eli Pariser (2021): The Filter Bubble. What the Internet is Hiding from You New York. The Penguin Press. 9-10.o.

köthető lehetséges veszélyeket is.²⁸³ A visszhangkamra-hatás az az - online és offline környezetben egyaránt megfigyelhető - jelenség, amikor az emberek hasonló gondolkodású közösségekbe csoportosulnak, ahol főként olyan információkkal és véleményekkel találkoznak, amelyek megerősítik saját előítéleteiket és nézeteiket. Kevesebb mint egy évtizeddel később megjelent könyvében Sunstein már arról ír, hogy a felhasználókat aktív közreműködésük nélkül is visszhangkamrákba kényszerítik a közösségi média algoritmusai, ami nagyban hozzájárul az ún. *digitális enklávék* létrejöttéhez.²⁸⁴ Az olyan csoportok pedig, akik tartósan nem kommunikálnak a sajátjuktól eltérő véleménnyel bíró közösségekkel, hamar elszigeteltté válnak az online térben, ami hosszútávon az online közösségek ideológiai vagy érdeklődési kör alapján történő fragmentálódásához, (kiber)balkanizációhoz vezet.²⁸⁵

Ehhez kapcsolódik az is, hogy az internet elterjedése torzította az emberek észlelését arról is, hogy bizonyos álláspontok és nézetek mennyire elterjedtek a valóságban. Míg az offline világban az egyén korlátozott számú, többnyire a közelében élő személyekkel való interakció alapján szerzett információt arról, hogyan gondolkodnak mások, az online világban ezek, a fizikai határok megszűnnek, és hozzáférés nyílik egy információban ugyan gazdag, de azok hitelességét és reprezentativitását tekintve szélsőségesen aránytalan környezethez. A közösségi médiában a marginális csoportok is látszólag széles körben elterjedtnek tűnhetnek, ha kellően aktív és zajos közösségeket alkotnak.²⁸⁶ Ez könnyen azt a téves benyomást keltheti, hogy, az extrém vélemények is széles körben elfogadottak és sokak által támogatottak. A valós életben nehéz olyan emberekkel találkozni, akik hisznek abban, hogy a Föld lapos, míg az interneten, a közösségi média több milliárd aktív felhasználója között vannak, akik ebben hisznek, és most könnyen megtalálhatják és kapcsolódhatnak egymáshoz.²⁸⁷ Az ilyen aktív és hangos véleményközösségek legitimációt biztosítanak egy személy hitének, ami még inkább

²⁸³ Sunstein, Cass R. (2001): Republic.com Princeton NJ: Princeton University Press

²⁸⁴ Sunstein hangsúlyozza, hogy az ilyen visszhangkamrák és digitális enklávék elősegítheti a politikai és társadalmi csoportpolarizációt, ami a vélemények és attitűdök még szélsőségesebbé válásával, súlyos akadályt lehet az érdemi, konstruktív párbeszédnek és a nézetek ütköztetésének. Lásd: Sunstein, Cass (2009): Republic.com 2.0. New Jersey, Princeton University Press. 11.o.; Sunstein (2017): #Republic: Divided Democracy in the Age of Social Media. Princeton, Princeton University Press.

²⁸⁵ Az ilyen online környezetben a különböző csoportok saját információs univerzumokban működnek, eltérő tényekkel és narratívákkal. Közöttük kevés az interakció, nem hallják meg egymás álláspontját. A csoportidentitás megerősödik, az ellenvélemények elutasítása fokozódik. Megnő az előítéletek, sztereotípiák és szélsőséges nézetek terjedésének esélye. Általánosságban elmondható, hogy a társadalmi kohézió gyengül. L. ásd: Sunstein, Cass (2009): Republic.com 2.0. New Jersey, Princeton University Press. 79–80.o.

²⁸⁶ Leviston, Z., I. Walker, and S. Morwinski (2013): Your opinion on climate change might not be as common as you think. Nature Climate Change, 3:334–337.o.

²⁸⁷ A Föld lakosságának megközelítőleg 62%, kb. 5 milliárd ember használt már valamilyen típusú közösségi médiát (csak 2022-ben 266 millió új felhasználó csatlakozott először közösségi platformra. Bővebben: Meltwater & We Are Social (2024): Digital 2024 Global Overview Report. Elérhető: <https://datareportal.com/reports/digital-2024-global-overview-report>

nehezebbé teszi az elgondolásuk megváltoztatását. Ezt a folyamatot egy sor az emberi Információfeldolgozást érintő heurisztikai és kognitív torzítás is erősíti, melyek közül a legfontosabbak a (1) hamis konszenzus hatás (2) a megerősítése torzítás és a kognitív disszonancia (3) valamint a hitbeli meggyőződés állandósága (belief perseverance).

(1): A "hamis konszenzus hatás" (false consensus effect) arra a pszichológiai jelenségre utal, amely során az emberek hajlamosak feltételezni, hogy saját nézeteik, hiedelmeik és értékrendjük általánosabbak és elterjedtebbek, mint valójában. (2) A megerősítési torzítás okán az emberek hajlamosak azokat az információkat keresni, előtérbe helyezni és mélyebben az emlékezetükbe vésni, amelyek összhangban vannak és megerősítik előzetes meggyőződéseiket, nézeteiket és értékrendszerüket.²⁸⁸ Amikor olyan információkkal találkoznak, amelyek ellentmondanak a meglévő hiedelmeiknek, kognitív disszonanciát élhetnek át.²⁸⁹ Ennek feloldására hajlamosak lehetnek figyelmen kívül hagyni vagy elutasítani a tényeket, és ragaszkodni a saját meggyőződéseikhez, még akkor is, ha azok nem állják ki a valóság próbáját. (3): Megfigyelhető az is, hogy emberek gyakran akkor sem változtatják meg kialakult hiedelmeiket, ha új, azoknak ellentmondó információkkal szembesülnek (belief perseverance). Ez azzal magyarázható, hogy ítéletalkotásuk során ok-okozati összefüggéseket feltételeznek a világ eseményei között, és az okozati következtetéseket mélyen elraktározzák az emlékezetükben. Ha később ki is derül, hogy a kialakult következtetésük alapjául szolgáló információ hamis volt, az abból eredő következtetés továbbra is befolyásolja a gondolkodásukat (az eredeti információ cáfolata, nem írja automatikusan felül az abból kialakított és már rögzült következtetést vagy világmagyarázatot.)²⁹⁰ Ebben az információfeldolgozás terén

²⁸⁸ Nickerson, Raymond S. (1998): Confirmation bias: A ubiquitous phenomenon in many guises. Review of General Psychology 2 (2), 175–220.o.

²⁸⁹ A kognitív disszonancia – melynek elméletét Leon Festinger amerikai szociálpszichológus dolgozta ki 1957-ben - egyfajta mentális feszültséget vagy kellemetlen érzést takar, amely akkor lép fel, amikor egy személy egyszerre tart fenn egymással összeegyeztethetetlen gondolatokat, attitűdöket vagy cselekedeteket. Festinger elmélete alapján az emberek folyton törekednek a belső kognitív egyensúly elérésére. Amikor információk, vélemények vagy cselekedetek ütköznek egymással, feszültség keletkezik, és az egyén igyekszik csökkenteni ezt a feszültséget. A jelenségről először a *When Prophecy Fails* c. 1956-ban megjelent könyvében írt, Henry Riecken és Stanley Schachter társszerzőkkel. A könyv egy amerikai UFO-vallás csoportot követ nyomon, amely azt hirdette, hogy a világ hamarosan véget ér, és csak őket menti meg egy földönkívüli civilizáció. Festinger és társai beépültek a csoportba, és megfigyelték a tagok viselkedését, különösen azt, hogyan reagáltak, amikor a jóslatuk nem teljesült. A csoport tagjai nem vetették el hiedelmeiket, hanem új magyarázatokat kerestek, hogy fenntarthassák meggyőződéseiket, ami példázza a kognitív disszonancia csökkentésére irányuló törekvéseket. A témáról lásd bővebben: Festinger, L., Riecken, H.W., & Schachter, S. (1956): *When Prophecy Fails*. University of Minnesota Press; Festinger, L. (1957): *A Theory of Cognitive Dissonance*. Stanford University Press.

²⁹⁰ E jelenség mutatja, hogy nem elég egyszerűen cáfolni a hamis információkat. Új, hiteles magyarázatokat is kell adni az eseményekre, hogy az emberek le tudják cserélni a téves következtetéseiket. Például, ha valaki azt a téves információt hallja, hogy az oltások autizmust okozhatnak, melyre alapozva arra a következtetésre juthat, hogy az oltások veszélyesek, és nem akarja beoltatni a gyermekeit. Hiába szembesül később hiteles cáfolatokkal, amelyek egyértelműen bizonyítják, hogy az oltások nem okoznak autizmust, az eredeti hiedelme megmaradhat, és továbbra

pszichológiailag és technológiailag is kiszolgáltatott környezetben lépett be az emberiség a dezinformációktól terhes post truth (posztigazság) korszakba .

6.5 Az igazság hanyatlása és a tényeken túli valóság

A "post-truth" kifejezés egy olyan társadalmi és politikai környezetre utal, ahol az érzelmek és a személyes meggyőződések nagyobb szerepet játszanak a közvélemény formálásában, mint az objektív igazság. Ebben a környezetben az emberek döntéseit inkább az érzéseik és saját hiedelmeik vezérlik, semmint a valós, ellenőrizhető tények. Krekó Péter a "Tömegparanoia 2.0" című könyvében ezt úgy foglalja össze, hogy: *A post-truth mágikus világában elmosódik a tények és a vélemények közti különbség, és ugyanarról a jelenségről több, egymást kizáró „tény” is létezhet.*²⁹¹

A kifejezés a 2010-es évek második felében vált széles körben ismertté, különösen a 2016-os Brexit népszavazás és az amerikai elnökválasztási kampányok kapcsán.²⁹²Ezek során a politikai szereplők a korábbiaknál is gyakrabban terjesztettek érzelmekre ható, leegyszerűsített, félrevezető vagy hamis információkat, amelyeket a hallgatóság sokszor feltétel nélkül elfogadott, függetlenül azok valóságtartalmától. A post-truth korszakra jellemző politikai kommunikációt jól szemlélteti Kellyanne Conway, Donald Trump tanácsadójának 2017. január 22-én a Meet the Press televíziós műsorban adott interjúja is. Ebben Conway megvédte Sean Spicer sajtótitkár azon állítását, hogy Trump beiktatásán rekordméretű tömeg volt, annak ellenére, hogy az eseményről készült fotók és videós beszámolók ennek jól láthatóan ellentmondtak. Conway azt mondta, hogy Spicer csak alternatív tényeket (alternative facts)

is befolyásolhatja a döntéseit. Ahhoz, hogy megváltoztassa a véleményét, nem elég pusztán cáfolni a hamis információt. Szüksége van egy új, hiteles magyarázatra is, amely meggyőzően bemutatja, hogy miért biztonságosak és fontosak az oltások (az oltások szigorú biztonsági és hatékonysági teszteken mennek át, az oltást elutasítók gyermekei sokkal nagyobb eséllyel kaphatnak el súlyos betegségeket stb. Lásd: Brendan Nyhan, Jason Reifler (2015): Displacing Misinformation about Events: An Experimental Test of Causal Corrections, *Journal of Experimental Political Science*. Volume 2, Issue 1, 81-93.o.

²⁹¹ Krekó Péter (2021): Tömegparanoia 2.0 - Összeesküvés-elméletek, álhírek és dezinformáció. Budapest, Athenaeum Kiadó. 22.o.

²⁹² Jól érzékelteti, hogy ezekben az években mekkora figyelem összpontosult a jelenségre, az is, hogy az Oxford Dictionaries 2016-ban a post-truth kifejezést választotta az év szavának. A meghatározás szerint a post truth olyan helyzetekre utal, amikor *a tárgyilagosságot kevésbé hatnak a közvéleményre, mint az érzelmeken, személyes hiten alapuló érvek* ("relating to or denoting circumstances in which objective facts are less influential in shaping

közölt, amikor azt állította, hogy ez volt a legnagyobb közönség, amely valaha is jelen volt egy beiktatáson.²⁹³

Kétségtelen, hogy az igazság hanyatlásának²⁹⁴, folyamata nem a 2010-es években indult meg. A történelem számos példát szolgál arra nézve, hogy különböző totalitárius rendszerek miként igyekeztek az igazságot relativizálni és a valóságot elferdíteni a hatalmuk biztosítás érdekében.²⁹⁵

A modern digitális környezetben mégis azért tárgyalandó e témakör, mert az előbbieken vázolt algoritmusokból eredő információs torzulás, illetve az emberre jellemző pszichológiai sajátosságok különösen sebezhetővé tették a modern társadalmakat. Erre az információfeldolgozással kapcsolatos sebezhetőségre vezethető vissza az is, hogy napjainkban ilyen gazdag táptalajra találtak a legkülönbélebb első hallásra a valóságtól és racionális belátástól teljesen elrugaszkodottnak tűnő összeesküvés elméletek (többek között azok a mélyállamról szóló konspirációk is, amelyek szerepet játszottak az Egyesült Államok Capitoliumának 2021. január 6-án bekövetkező ostromában).²⁹⁶

²⁹³ Chuck Todd, a műsor házigazdája, erre úgy reagált, hogy "az alternatív tények nem tények, hanem valótlanságok (alternative facts are not facts, they are falsehood). Conway megpróbálta elterelni a beszélgetést más témákra, mint például az egészségügyi biztosítás kérdése, anélkül, hogy közvetlenül válaszolt volna a hamis állításra vonatkozó kérdésekre. Lásd: Gajanan, Mahita (2017): Kellyanne Conway Defends White House's Falsehoods as 'Alternative Facts'. TIME. Elérhető: <https://time.com/4642689/kellyanne-conway-sean-spicer-donald-trump-alternative-facts/> (2017. 01. 12.)

²⁹⁴ Az igazság hanyatlása "truth decay" kifejezést Kavanaugh és Rich a RAND amerikai nonprofit politikai kutatóintézet tanácsadó megbízásából készített 2018-as jelentésében használta, arra hogy leírja az objektív tények szerepének csökkenését a közbeszédben és a politikai viták során. A jelenséget négy trenddel jellemezték: (1) A tényekbe vetett bizalom csökkenése (2) A vélemények és a tények közötti különbségek elmosódása (3) A személyes vélemény felértékelődése a tényekkel szemben (4) A korábban hitelesnek elismert információforrásokba vetett bizalom csökkenése Kavanaugh, Jennifer and Michael D. Rich (2018): Truth Decay: An Initial Exploration of the Diminishing Role of Facts and Analysis in American Public Life. Santa Monica, CA: RAND Corporation. 21-38.o. Elérhető: https://www.rand.org/pubs/research_reports/RR2314.html.

²⁹⁵ Az autoriter rendszerek sajátosságait kutató Hannah Arendt is mélyrehatóan vizsgálta az igazság jelentésének mibenlétét. Munkásságában különbséget tett a tényigazságok (factual truth) és az észigazságok (rational truth) között. Értelmezés szerint az észigazság filozófiai, tudományos, matematikai igazságokat takar, melyekhez az elme útján, gondolkodás révén lehetséges eljutni. Ezek állandóak és biztosak (logikai kényszerítő erővel bírnak) ellentétpárjuk a tévedés, a tudatlanság vagy az illúzió. A tényigazságok ezzel szemben egy adott helyzetre történésre, eseményre vonatkoznak, ami természetükből adódóan - különösen a politikai szférában - könnyen torzíthatóak, manipulálhatóak (például átalakulhatnak véleményekké, amikor politikai diskurzusok tárgyává válnak). A tényigazság ellentétpárját viszont nem a tévedés, vagy a tudatlanság adja, hanem a szándékos hazugság. A rendszerszinten gyakorolt szándékos hazugság a totalitárius berendezkedések sajátja. Lásd: Arendt, Hannah (1995): Igazság és politika. In: Múlt és jövő között. Nyolc gyakorlat a politikai gondolkodás terén. (ford: Módos Magdolna). Budapest, Osiris-Readers International, 233-251.o.

²⁹⁶ Még a 2016-os amerikai elnökválasztási kampány utolsó heteiben terjedt el a Pizzagate néven elhíresült konspirációs elmélet, mely szerint Hillary Clinton kampányfőnöke, John Podesta és más prominens demokrata politikusok gyermek szexkereskedelmi hálózatot üzemeltetnek egy washingtoni pizzéria (Comet Ping Pong) alagsorában. Bár erre semmilyen konkrét bizonyíték nem volt, a Wikileaks által korábban kiszivároztatott Clinton - Podesta privát e-mailek szövegeiben az elmélet hívei különböző rejtett utalásokat véltek felfedezni az állításaik alátámasztására. Erre az elméletre épült később a QAnon összeesküvés-elmélet, mely szerint egy rejtélyes "Q"

A helyzetet azonban tovább fokozza, hogy ebben a kiszolgáltatott közegben indultak meg azok a jól felépített, tudatos dezinformációs kampányok, amelyek a közvélemény befolyásolását tűzték ki céljukként. Jól illik ebbe az új információs környezetbe a modern dezinformáció természete, amely sok esetben nem is a nem meggyőzésre, hanem elbizonytalanításra törekszik.

297

6.6 Félretájékoztatás, dezinformáció

A 2016-os amerikai elnökválasztásban az álhírek és a dezinformációs kampányok kiemelt szerepet játszottak. Ebben az időszakban az orosz kormánnyal közvetlen kapcsolatba hozható trollfarmok nagy mennyiségben terjesztettek megosztó, és félrevezető tartalmakat a közösségi médiában, hogy befolyásolják a választások kimenetelét és mélyítsék a politikai polarizációt az USA-ban.²⁹⁸ Az erre szakosodott műhelyek közül a legismertebb az azóta váratlan repülőbalesetben elhunyt, ex-Wagner vezér Jevgenyij Prigozsин tulajdonában álló szentpétervári Internet Research Agency (IRA) volt. A hibrid hadviselés²⁹⁹ részét képező dezinformációs művelet kiterjedtségét mutatja, hogy a Facebook az amerikai szenátus igazságügyi bizottsága előtt elismerte, hogy az Oroszország által támogatott digitális tartalmak

nevű kormányzati hírszerző információkkal rendelkezik egy globális elit csoportról, akik pedofil hálózatokat üzemeltetnek. Donald Trump, az Amerikai Egyesült Államok 45. elnöke, azonban harcot indított ezen elit csoport ellen, és titokban küzd az ő globális befolyásuk csökkentéséért. Az elmélet hívei a 2020-as elnökválasztáson Trump vereségét úgy értelmezték, mint a szabadságharcosok bukását a korrupt háttérhatalom ellen. Ebben a feltűzött helyzetben kérdőjelezte meg Trump az elnökválasztás eredményét és vádolta csalással a demokratákat. A választást követő két hónapban egyre szélesebbé és hangosabbá váltak azok a visszhangkamrák, amik a korrupt és romlott demokrata politikusok felelősségét kiemelve a választási eredmények rendszerszerű meghamisítását hangoztatták. A leköszönő elnök január 6-ai beszédét követően a reményvesztett tömeg Capitoliumhoz vonult, áttörték a kordonokat, összecsaptak a rendőrökkel, betörték az épületbe, amit több órán át megszállva tartottak. Az ostromban számos QAnon követő is részt vett, a leghíresebb közülük Jake Angeli volt, akit a médiában csak "QAnon sámán"-ként emlegettek.

²⁹⁷ Krekó Péter, Molnár Csaba (2023): Tényrelativizmus és a hírforrásokkal szembeni elbizonytalanodás a magyar közvéleményben. Lakmusz-HDMO. 5.o. Elérhető: https://politicalcapital.hu/pc-admin/source/documents/hdmo_pc_tanulmany_2_tenyrelativizmus_kozvelemenye_20231130.pdf (2023.11.30.)

²⁹⁸ Az orosz beavatkozást az amerikai hírszerző szervek is megerősítették, kiemelve, hogy Oroszország célja a választási folyamat befolyásolása és az amerikai politikai rendszer meggyengítése volt. Egy az Egyesült Államok szenátusának hírszerzési bizottsága által közzétett több száz oldalas jelentés megállapítja, hogy bár Moszkva korábban is szivárogtatott ki politikailag érzékeny információkat, illetve 2014 óta egyre nagyobb mennyiségben tett elérhetővé harmadik fél közreműködésével (pl.: WikiLeaks) illegálisan szerzett iratanyagokat, a 2016-os amerikai elnökválasztás során tanúsított aktivitása egyértelműen szintlépésnek tekinthető. Lásd bővebben: U.S. Government Publishing Office (2019): Report of the Select Committee on Intelligence United States Senate on Russian Active Measures Campaigns and Interference in the 2016 U.S. Election. Senate Report, II. Volume, 116–290o. Elérhető: www.intelligence.senate.gov/publications/report-select-committee-intelligence-united-states-senate-russian-active-measures

²⁹⁹ A hibrid hadviselés olyan katonai stratégia, amely egyesíti a hagyományos katonai műveleteket a nem hagyományos módszerekkel, mint a kibertámadások, dezinformációs kampányok, gazdasági nyomásgyakorlás, valamint politikai destabilizáció. Ezen módszerek célja a közvélemény befolyásolása, a politikai stabilitás aláásása és a döntéshozatali folyamatok megzavarása.

a 2016-os amerikai elnökválasztás alatt és után körülbelül 126 millió amerikait felhasználót értek el a közösségi platformjukon. Ebben az időszakban az Oroszországhoz köthető oldalak több mint 80 ezer bejegyzést hoztak létre, amelyek közvetlenül 29 millió amerikai felhasználóhoz jutottak el. Azonban, a másodlagos megosztásokat és interakciókat is számításba véve a bejegyzések - eredeti címzettekén túl, a felhasználók barátainak és követőinek hálózatába terjeszkedve - valójában közel ötször annyi emberhez jutottak el. A Twitter esetében még kiterjedtebb kampányról volt szó, melyben 50 ezer automatizált felhasználói fiók (bot) hozzávetőlegesen 1,4 millió tweetet tett közzé, közel 288 millió interakciót generálva.³⁰⁰

Bár az IRA amerikai tevékenysége figyelemre méltó példa, a hibrid hadviselés keretein belül folytatott félretájékoztatási kampányok nem korlátozódnak egyetlen országra vagy szereplőre. A világ különböző állami és nem állami szereplői hasonló tevékenységet folytatnak, politikai, ideológiai vagy gazdasági haszonszerzés céljából.³⁰¹ Például az elmúlt években az Európai Unióra összpontosuló dezinformációs hibrid hadműveletek olyan világnézeteket és politikai erőket karoltak fel - és törekedtek felerősíteni - a propagandájukkal, amelyek az európai integráció és transzatlanti együttműködés ellen foglalnak állást, és a status quo felborítását célozzák. Igyekeztek kiélezni a belső megosztottságot a célországokban és felnagyítani a szélsőséges hangokat olyan érzékeny társadalmi ügyek kapcsán mint a migráció, szexuális kisebbségek jogai, eutanázia, terhességmegszakítás vagy drogliberalizáció. Az oltásokkal kapcsolatos álhíreket és összeesküvés-elméleteket felkapva igyekeztek aláásni a bizalmat a tudományos intézményekben és szakértőkben, ezzel gyengítve a járványkezelési erőfeszítéseket a nyugati országokban. Geopolitikai válsághelyzeteket, mint a brexit, a katalán függetlenségi népszavazás, vagy egyes nagyszabású terrortámadások utóhatásait meglovagolva terjesztettek zavar keltő álhíreket, fokozva a bizonytalanság és a széthúzás légkörét.

³⁰⁰ Jack Dorsey a twitter akkori vezérigazgatója utólag vallotta csak be, hogy ez idő tájt milyen nagy számban voltak jelen az Internet Research Agencyhez köthető hamis fiókok a felületükön. Az elnökválasztásról folyamatosan tweetelő automatizált fiókok (botok) számát kezdetben 36.000 becsülték, amely számot később korrigáltak 50.000-re.

³⁰¹ A nemzetközi politikára gyakorolt hatását mutatja, hogy a Világgazdasági Fórum a 2024-es Global Risk Report jelentése az elkövetkező két év kiemelt és már rövid-távon is meghatározó kockázatait a következő módon rangsorolja: (1) dezinformáció és félretájékoztatás globális elterjedése, (2) extrém időjárási események, (3) társadalmi polarizáció, (4) kiberbiztonsági fenyegetések, (5) államközi fegyveres konfliktusok, (6) gazdasági lehetőségek hiánya, (7) infláció, (8) kényszerű migráció, (9) gazdasági visszaesés, és (10) környezetszennyezés. Lásd: World Economic Forum (2024): The Global Risk Report 2024 - 19th Edition. 14-37o.

6.7 A dezinformáció szolgálatában - Mikro-célzás, állhírek és botok

Az Európai Bizottság meghatározása szerint a dezinformáció *„olyan igazolhatóan hamis vagy félrevezető információ, amelyet gazdasági haszonszerzés vagy szándékos megtévesztés céljából hoznak létre, hoznak nyilvánosságra és terjesztenek, és amely kárt okozhat a közérdeknek”*.³⁰²

A dezinformációk tartalmáról általánosságban elmondható, hogy gyakran közölnek leegyszerűsített narratívákat, amelyek az összetett kérdések megértése helyett arra ösztönzik a célszemélyeket, hogy figyelmen kívül hagyják az adott ügy bonyolultabb, árnyaltabb részleteit. Emellett előszeretettel használnak olyan erős érzelmi elemeket, mint a félelem, a düh és a bizonytalanság, amely nem csak nagyobb elérést és gyorsabb terjedést biztosít a figyelemgazdaság korában, de egyúttal erőteljesebben is képes befolyásolni az az emberi ítélőképességet, ezáltal fogékonyabbá téve a fogyasztót a manipulációra.³⁰³³⁰⁴ A dezinformációs kampányok célja sokrétű; a nyilvánosság félrevezetésétől kezdve, egy adott témával kapcsolatos többségi vélemény kialakításán keresztül, a szabad és tisztességes választások befolyásolásáig sok minden lehet. Ennek elérésére különböző megoldások állnak rendelkezésre:

A kereskedelmi gyakorlatokban elterjedt online hirdetési módszerek, különösen a mikro-célzás (micro-targeting) technikák nem csak a termékek reklámozásában hasznosak, hanem a félrevezető információk terjesztésében is. A mikro-célzás lényege, hogy a hirdető a

³⁰² Fontos a dezinformációtól elkülöníteni félretájékoztatás fogalmát, amely olyan hamis vagy félrevezető tartalom terjesztését jelenti, amelyet káros szándék nélkül osztanak meg (ám hatásai ettől még kedvezőtlenek) Szintén nem tartoznak e kategóriába a jelentéstételi hibák, a satírák és paródiák, illetve az egyértelműen azonosítható módon pártokhoz köthető hírek és kommentárok. Lásd: Európai Bizottság (2018): Európai megközelítés az online félretájékoztatás kezelésére. COM(2018) 236 final, 2.1. pont.

³⁰³. A negatív hírek erőteljesebb érzelmi reakciókat váltanak ki, mint a pozitív vagy semleges hírek. Ennek oka, hogy az emberi agy hajlamosabb nagyobb figyelmet fordítani a negatív információkra, ami az evolúciós múltból ered, ahol a negatív információk, mint a veszély vagy fenyegetés, fontos túlélési jelzések voltak (az olyan erős érzelmek, mint a düh vagy a félelem pedig, mélyebb emlékezetei rögzülést is eredményezhetnek). Bővebben: Robertson, C. E., Pröllochs, N., Schwarzenegger, K., Pärnamets, P., Van Bavel, J. J., & Feuerriegel, S. (2023). Negativity drives online news consumption. *Nature human behaviour*, 7(5), 812–822.o.

³⁰⁴ Látható tehát, hogy a dezinformáció terjedésének gyorsaságában szerepet játszik előzőekben említett információfeldolgozási torzulás mindkét formája. (1) Technológiai torzítás: A személyre szabott tartalomszűrő algoritmusok a magas aktivitást kiváltó tartalmakat helyezik előtérbe, és információs buborékokat hoznak létre (2) Pszichológiai torzítás: az emberek erőteljesebben reagálnak az érzelmekre ható negatív üzenetekre, elfogultak előzetes tudásukat illetően, túlbecsülik saját ismeretüket egy adott témában, illetve a figyelemküszöb csökkenése okán a rövid és könnyen emészthető tartalmakat részesíti előnyben. Ennek eredményeképpen a szenzációhajás dezinformációs narratívák gyakran vírusként terjednek, míg a pontos és tényszerűen ellenőrzött (terjedelmesebb, szárazabb, érzelmentes, nehezebben befogadható) információknak hosszabb időbe telik, mire teret nyernek. Ráadásul az is súlyosbítja a problémát, hogy a közösségi médiákban a felhasználó fogyasztó és tartalom-előállító is egyben, így a megosztások, kedvelések és hozzászólások által könnyen maga is dezinformációs kampányok propagálójává válik.

felhasználókról gyűjtött személyes adatok alapján pontosan meg tudják határozni, kik tartoznak a célcsoportjukba, így az üzeneteiket aprólékosan személyre szabhatják, növelve ezzel azok hatékonyságát.³⁰⁵ A politikai mikrocélzás legveszélyesebb formái közé azok tartoznak, melyek kifinomult pszichometriaival profilalkotási módszereket alkalmaznak az emberek személyes sebezhetőségeit használják ki.³⁰⁶ A félretájékoztatási műveletek esetében az ilyen mikro-célzás abban segíthet, hogy a megtévesztő vagy hamis információkat célzottan juttassák el azokhoz a csoportokhoz, amelyek a legfogékonyabbak rájuk. Például, ha valaki összeesküvés-elméleteket akar terjeszteni, a mikro-célzással azonosíthatja azokat a felhasználókat, akik korábban hasonló tartalmakkal léptek interakcióba, majd célzottan nekik küldheti az üzeneteket.

A közvélemény befolyásolásának szándékával vagy a sok esetben egyszerű anyagi haszon realizálásának reményében³⁰⁷ terjesztett álhírek (fake news) is fontos eszközei lehetnek a dezinformációnak. Az ilyen tartalmaknál a szenzációhajhász cím mellett előszeretettel használnak internetes mém vagy infografika formátumokat, amivel az összetett kérdéseket leegyszerűsítve és vizuálisan is vonzó, könnyen befogadható módon közvetítik.³⁰⁸ Itt is érvényes a szabály, hogy az érzelmekre ható álhírek különösen hatékonyak, mivel könnyen felkeltik az olvasók érdeklődését, egyúttal ösztönzik a megosztásokat. A megtévesztés érdekében ezek forrásaiként sok esetben olyan weboldalak szolgálnak, amik felépítésükben elrendezésükben, és domain nevükben is hasonlóak lehetnek, mint a legitim, ismert hírügynökségek.³⁰⁹

³⁰⁵ Talán a legismertebb példa erre, a Cambridge Analytica politikai tanácsadó cég esete, amely mintegy 87 millió Facebook felhasználó adataihoz férte hozzá a felhasználók beleegyezése nélkül. A cég a 2016-os amerikai választásokon ezeket az adatokat pszichológiai profilozásra és a szavazók politikai hiedelmekkel való célzott megszólítására használta. Bővebben: Isaak, J., & Hanna, M. (2018): User Data Privacy: Facebook, Cambridge Analytica, and Privacy Protection. *Computer*. 51, 56-59.o. <https://doi.org/10.1109/MC.2018.3191268>.

³⁰⁶ . Európai Bizottság (2021): A dezinformáció elleni gyakorlati kódex megerősítésére vonatkozó iránymutatás. COM(2021) 262 final, 11-12o

³⁰⁷ Az Észak-Macedóniai Veles városában például több tucat, főleg fiatalok által működtetett weboldal specializálódott ennek kiaknázására. A városban működő weboldalak egyszerű, de annál hatékonyabb gazdasági modellje az álhírek gyártásán és terjesztésén, főként amerikai olvasók célzott bevonzásán, majd a hirdetések megjelenítése révén a nyereség realizálásán alapult (a Google AdSense szolgáltatásának segítségével). Kihhasználva az amerikai politikai polarizációt és az online hírfogyasztási szokásokat a Veles-i weboldalak tipikusan amerikai politikával kapcsolatos tartalmakat gyártottak - például Hillary Clinton állítólagos korrupciós bűncselekményeiről vagy Barack Obama „eltitkolt, valódi” születési helyéről. Ezek a weboldalak aztán a közösségi médiában, különösen a Facebookon keresztül, terjesztették ezeket a tartalmakat. Bővebben: Hughes, H., & Waismel-Manor, I. (2020): The Macedonian Fake News Industry and the 2016 US Election. *PS: Political Science & Politics*, 54, 19 – 23.o. <https://doi.org/10.1017/S1049096520000992>.

³⁰⁸ Veszelszki Ágnes (2017): Az álhírek extra- és intralingvális jellemzői. *Századvég* 84: 51–82.o.

³⁰⁹ Ilyenek például azok a kérészetű, Magyarországon is rendszerint a választásokat megelőző hónapokban aktivizálódó, tisztán pártpolitikai érdekek mentén működő, impresszum nélküli online kiadványok, melyek nincsenek bejegyezve az NMHH-nál, nem minősülnek sajtóterméknek és amelyekkel szemben sajtóhelyreigazítási per sem indítható. Ezek a weboldalak néhány hónapon keresztül – főleg a kampányidőszakban – osztanak meg cikkeket, melyek sok esetben erős keretezéssel, torzítással közöltek – fő célok között a

A dezinformáció terjedésének fokozását nagyban segítheti a hamis profilokat irányító botok bevetése. Ezek olyan automatizált szoftverek, amiket olyan egyszerű online művelet elvégzésére hoztak létre, mint a bejegyzések megosztása, komment kedvelése, hashtagek felfuttatására stb. Ezek egyik altípusát, a politikai botokat gyakran használják arra, hogy a választási kampányokban egy jelölt vagy politikai párt üzeneteit terjesszék, de akár agresszív kampányolásra is programozhatják mely során úgy képesek manipulálni, egyúttal torzítani a nyilvános diskurzust, hogy propagandát terjesztenek, újságírókat támadnak, vagy éppen politikai vezetőket járatnak le.³¹⁰ Samuel Woolley a politikai botokkal kapcsolatban megjegyzi, hogy elterjedt félreértés, hogy azokat úgy tervezték, hogy a felhasználókkal kommunikáljanak és befolyásolják azok nézeteiket. Ezek elsődleges feladata valójában az, hogy a közösségi média platformok algoritmusait manipulálják, mivel ezek döntenek el, hogy milyen tartalmak jelennek meg hangsúlyosan a felhasználóknak. A botok oly módon igyekeznek befolyásolni ezt a folyamatot, hogy mesterségesen felduzzasztják bizonyos tartalmak népszerűségét, például sok lájkkal, megosztással, hozzászólással, így segítve hozzá azokat, hogy a lehető legtöbb felhasználóhoz jussanak el.³¹¹ Fontos, hogy minél többször minél több helyen jelenjenek meg az ilyen tartalmak, hiszen. A dezinformáció módszertanához közé tartozik az ismétlés általi megerősítés, ahol a hamis információkat gyakran és különböző csatornákon ismétlik, hogy legitimitást és hitelességet sugalljanak.³¹²

dezinformáció, karaktergyilkosság és lejáratás áll. A választások végével ezen weboldalak aktivitása, és közösségi média hirdetései is elapadnak. Lásd bővebben: Ambrus Balázs (2023): Megtévesztő híroldalakkal építik magukat a pártok. Index. Elérhető: <https://index.hu/belfold/2023/11/07/impreszum-lapok-kommunikacio-helyi-hirek-propaganda/>

³¹⁰Howard, P. N., Woolley, S., & Calo, R. (2018): Algorithms, bots, and political communication in the US 2016 election: The challenge of automated political communication for election law and administration. *Journal of information technology & politics*, 15(2), 81-93.o.

³¹¹Wooley, Samuel C. (2020): Bots and Computational Propaganda: Automation for Communication and Control. In *Social Media and Democracy: State of the Field*. (Szerk:) Persily, Nathaniel and Tucker, Joshua A. Cambridge, Cambridge University Press, 89–110.o.

³¹² A botok hasznosságát szemlélteti a Kínai Kommunista Párthoz (KKP) köthető "Spamouflage" online dezinformációs művelet, amely aktuális nemzetközi konfliktusokat, globális eseményeket felhasználva igyekszik társadalmi és politikai polarizációt előidézni, valamint Amerika-ellenes vagy a KKP-t támogató narratívákat terjeszteni. Stratégiájuk, hogy fényképek, videók és hírcikkek képernyőképeit használják (más esetekben GenMI által létrehozott képeket) melyek látványos, figyelemfelkeltő, szarkasztikus grafikai illusztrációkként szolgálnak fő narratívák képi közléséhez. A kampány gerincét azonban nem a tartalomgyártás, hanem azok a botfiókok alkotják, melyek egy viszonylag kisebb csoportja a tartalomközzétételére, míg egy nagyobb csoport, a bejegyzések kedvelésével, megosztásával és kommentálásával foglalkozik (annak érdekében, hogy felnagyítsák a bejegyzések láthatóságát és hatását). Ez a kampány több mint 50 különböző platformon és fórumon volt aktív, beleértve a Facebookot, Instagramot, TikTOKot, YouTube-ot és a Twittert (most már "X" néven ismert). A Meta jelentése szerint több mint 7,700 Facebook-fiókot távolítottak el, amelyek részt vettek ebben a kiterjedt online „spam” műveletben Lásd: Josh Taylor (2023): Meta closes nearly 9,000 Facebook and Instagram accounts linked to Chinese ‘Spamouflage’ foreign influence campaign. *The Guardian*. Elérhető: <https://www.theguardian.com/australia-news/2023/aug/30/meta-facebook-instagram-shuts-down-spamouflage-network-china-foreign-influence> (2023. 08. 30.)

6.8 GenMI a dezinformáció szolgálatában

Az előző évtizedben indított dezinformációs műveletek még több korlátba is ütköztek. Ezek közül kiemelendő a rendelkezésre álló erőforrások korlátozottsága, az üzenetek és közvetített tartalmak minőségének alacsony színvonala³¹³, valamint az akciók könnyű észlelhetősége. Mivel ezek az alacsony minőségű üzenetek, könnyen észlelhetők és felismerhetőek voltak, korlátozott volt a kampányok hatékonysága (a célközönség gyakrabban utasította el a hamis információkat).

A Nagy Nyelvi Modellek megjelenésével és szabadon hozzáférhetővé válásával ezen korlátok egy jelentős része lebontható. Az LLM segítségével könnyen, gyorsan, minimális anyagi költségekkel és minimális emberi erőforrás felhasználásával lehet nagy mennyiségű, hitelesnek tűnő dezinformációt előállítani. A felhasználók egyéni preferenciáihoz igazodó célzott dezinformáció is pontosabb lehet segítségével. Több kutatás is arra a következtetésre jutott, hogy a legújabb Nagy Nyelvi Modellek már kreatívabbak³¹⁴ és meggyőzőbbek³¹⁵ a szöveges tartalomgyártásban, mint az átlagemberek.³¹⁶ Bai és társai például egy 2023-as tanulmányukban számoltak be három különböző kísérletről, melyet 2022 októbere és 2022 decembere között végeztek, összesen 4836 résztvevővel. (1203, 2023, 1610 fő). A kísérletek során különféle politikai üzenetekkel próbálták megváltoztatni a résztvevők nézeteit olyan témákban, mint a

³¹³ Például visszatérő jelenség volt, hogy a dezinformációs kampányok sok esetben angol nyelvterületű országot céloznak meg, de a szöveges tartalmak előállítására nem angol anyanyelvű, alacsony képzettségű, olcsó munkaerőt alkalmaztak, akiknek a helyesírási és nyelvtani hibái hamar leleplezték az üzenetek kétes eredetét.

³¹⁴ Hubert és társai egy 2024-es tanulmányukban emberi résztvevők (N=151) és a GPT-4 GenMI válaszait hasonlították össze a következő három divergens gondolkodási feladat megoldásában (1): Alternatív felhasználási asszociáció (Alternative Uses Task, AUT), ahol résztvevőknek egy hétköznapi tárgy, újszerű kreatív felhasználási módjait kellett kitalálniuk (2): Lehetséges következmények feladat (Consequences Task, CT), mely során résztvevőknek képzeletbeli forgatókönyvek lehetséges következményeit kellett megfogalmazniuk (3): Divergens asszociációs feladat (DAT): A résztvevőknek 10 olyan főnevet kellett megadniuk, amelyek kinézetük, tartalmuk, funkciójuk stb. tekintetében a lehető leginkább különbözőek. Mindhárom vizsgált feladatra igaz volt, hogy az LLM válaszai egyrészt eredetibbek voltak, (azaz egyedibbek, szokatlanabbak, "out-of-the-box" jellegűek) és kidolgozottabbak is (vagyis részletesebbek, hosszabbak, alaposabbak) mint az emberi résztvevők válaszai. Az LLM-ek esetében feltűnő volt azonban, hogy nagyobb gyakorisággal használtak újra és újra ismétlődő szavakat válaszaik generálása során. Bővebben: Hubert, K. F., Awa, K. N., & Zabelina, D. L. (2024). The current state of artificial intelligence generative language models is more creative than humans on divergent thinking tasks. *Scientific reports*, 14(1), 3440. <https://doi.org/10.1038/s41598-024-53303-w>

³¹⁵ Goldstein A. Josh, Jason Chao, Shelby Grossman, Alex Stamos, Michael Tomz (2024): How persuasive is AI-generated propaganda? *PNAS Nexus*. Volume 3, No 2, 1-7o. ; Vykopal, I., Pikuliak, M., Srba, I., Móro, R., Macko, D., & Bielíková, M. (2023). Disinformation Capabilities of Large Language Models. *ArXiv*, abs/2311.08838. <https://doi.org/10.48550/arXiv.2311.08838>.

³¹⁶ Sőt egy 2024-es tanulmány megállapítása szerint a GPT 3.5 LLM már nem csak kreatívabb, de humorosabb is (200 a kísérletben résztvevő emberi bíráló értékelése alapján). Lásd: Gorenz D, Schwarz N (2024): How funny is ChatGPT? A comparison of human- and A.I.-produced jokes. *PLoS ONE* 19(7), 2-10.o. <https://doi.org/10.1371/journal.pone.0305364>

dohányzással kapcsolatos szabályozások szigorítása, fegyverviselési jogok korlátozása vagy a széndioxid adó bevezetése. A kísérletben három üzenetkategóriát használtak: (1) MI által generált érvek (AI Condition), (2) természetes személyek által létrehozott üzenetek (human condition), (3) öt MI által írt szöveg, melyek közül emberi döntés alapján került kiválasztásra egy (human in the loop condition). A résztvevők mindegyik esetben véletlen besorolás alapján kapták a politika tartalmú üzeneteket, így azok emberi vagy szintetikus jellegéről nem volt tudomásuk. Mindhárom kísérletben az MI által generált üzenetek következetesen meggyőzőnek bizonyultak a résztvevők számára (minimum 2-4 pontos változás a 101 pontos összesített attitűdskálán). Bár a különböző kategóriák közötti a meggyőzés hatékonysága statisztikailag elhanyagolható volt, a résztvevők értékelése szerint az MI által létrehozott szövegek inkább tűntek tényszerűnek, következetesnek és jól követhető racionális érvekből állónak.³¹⁷

Az OpenAI egy jelentésében arról írt, hogy az orosz Doppelgänger nevű csoport a ChatGPT-t használta fel olyan álhírek terjesztésére, amelyek célja az Ukrajna iránti szolidaritás és bizalom aláásása volt. A csoport több nyelven terjesztett olyan hamis információkat (jól felépített, és megfogalmazott Facebook-bejegyzések, cikkek, hozzászólások formájában) amelyek pro-orosz és ukrán ellenes üzeneteket közvetítettek.³¹⁸ A modern dezinformációs műveletek azonban már nem csak szöveges üzenetekre korlátozódnak, a hatékonyság fokozása érdekében kifinomult multimédiás tartalmakra is építenek

6.9 Szintetikus hang és videótartalmak

A vizuális tartalmak előre meghatározott céllal történő szándékos manipulációja nem új keletű jelenség, ilyen gyakorlatok már a fotózás és a filmkészítés kezdetei óta léteznek.³¹⁹ Korábban

³¹⁷ Ezzel szemben az emberek által írt üzenetek inkább személyes narratívákra és történeti elbeszélésekre épültek (pl: tapasztalatok, élmények, képi-leíró elemek használata stb.). A tanulmány azt a következtetést állapítja meg, hogy a GPT-3as modell által generált kimenetek, csakúgy mint a profi marketingesek által létrehozott politikai üzenetek, képesek arra, hogy érdemben befolyásolják az azt olvasó, fogyasztó személyek előzetesen kialakult nézetét. Lásd: Bai, Hui, Jan G. Voelkel, Johannes C. Eichstaedt, and Robb Willer. (2023): Artificial Intelligence Can Persuade Humans on Political Issues.” OSF Preprints. <https://doi.org/10.31219/osf.io/stakv>

³¹⁸ Például hamis híroldalakat hoztak és különböző hírességek nevével fiktív idézeteket generáltak melyeken keresztül, olyan hamis információkat közvetítettek, mint hogy a Nyugat Ukrajnának nyújtott támogatását Izraelre irányítják át.

³¹⁹ Már a Szovjetunióban is előszeretettel fordultak fényképek manipulációjához (retusálásához) a történelem újraírása érdekében. Ez különösen jellemző volt a sztálini időszakra, amikor a politikai tisztogatások során sok korábbi szövetségest és közismert személyt eltávolítottak a hivatalos történelemből. Az egyik leghíresebb példa erre, amikor Lev Trotszkijt, aki korábban Sztálin egyik legfőbb riválisa volt, eltávolították a hivatalos fotókról és

e manipulációs technikák alkalmazásához technikai tudásra és szakmai tapasztalatra volt szükség, mely jelentős anyagi és humánerőforrást igényelt. Napjainkra azonban jelentős változások mentek végbe ezen a területen, a modern, könnyen hozzáférhető, felhasználóbarát digitális technológiák lehetővé teszik a fotók és videók - magas minőségű - szerkesztését és manipulálását akár néhány egyszerű kattintással. Ezzel a vizuális tartalmak hamisítása olyan szintre jutott, ami korábban elképzelhetetlen volt. Az amatőr és professzionális felhasználók számára egyaránt elérhetővé váltak azok az eszközök, amelyek lehetővé teszik a képek és mozgóképek meggyőző meghamisítását.³²⁰

A GenMI modellek segítségével létrehozott vagy manipulált audiovizuális tartalmak (deepfake tartalmak)³²¹ járványszerű elterjedését valójában évtizedes technológiai folyamat előzte meg, melynek fontos állomását képezte az ún. antagonisztikus hálózatok (Generative Adversarial Networks, GAN) 2014-es megjelenése. A GAN a gépi tanulás innovatív területét képezi, ahol az új tartalom létrehozása két különálló neurális hálózat együttműködésén alapul; Ezek közül a *generátor hálózat* (generator) felelős az új adatok generálásáért, célja, hogy olyan adatokat hozzon létre, amelyek valóságosnak tűnnek (audio, kép, videó). A *diszkriminátor hálózat* (discriminator) pedig értékeli, hogy a generátor által létrehozott adatok valóságosak-e vagy sem, célja, hogy különbséget tegyen a valódi és a mesterséges adatok között. A GAN működése során a generátor és a diszkriminátor egyfajta (antagonisztikus) versenyben van egymással, ahol a generatív hálózat egyre jobban igyekszik utánozni a valóságot, hogy becsapja a diszkriminatív hálózatot, amely viszont egyre hatékonyabban próbálja felismerni és kiszűrni a generált adatokat. Ez a folyamatos versengés hozzájárul a generatív modell szüntelen fejlődéséhez és finomításához, így a generált adatok egyre nehezebben különböztethetők meg a valós adatoktól. Ennek köszönhetően a GAN-ok számos területen váltak alkalmazhatóvá, például

dokumentumokból, miután kegyvesztett lett. Az egyik legismertebb példa erre, amikor Lev Trotszkijt eltávolították a hivatalos fotókról röviddel Sztálin hatalomra kerülése után.

³²⁰ Guld Ádám (2023): A deepfake és CGI-technológia az influencer marketing szolgálatában: így formálják át a digitális karakterek az ismertségipar működését. In: Aczél, Petra; Veszelszki, Ágnes (szerk.) Deepfake: a valótlan valóság. Budapest, Gondolat Kiadó. 189-190.o.

³²¹ Olyan digitális médiatartalom (elsősorban audiovizuális formátumú, de lehet például kép, animáció vagy 3D-modell), amelyet részben vagy egészben mesterséges MI technológiák felhasználásával hoztak létre, illetve olyan már létező tartalom, amelyet MI-alapú eszközökkel módosítottak vagy manipuláltak. A rendelet 3. cikk (60) fogalom meghatározásában: az MI által generált vagy manipulált kép, audio- vagy videotartalom, amely hasonlít létező személyekre, tárgyakra, helyekre, entitásokra vagy eseményekre, és amely egy személy számára megtévesztő módon autentikusnak vagy valóságosnak tűnne.

képfeldolgozásban, művészeti alkotások generálásában, vagy akár realisztikus szimulációk létrehozásában is.³²²

Több kutatás is alátámasztja, hogy az MI által mesterségesen előállított vagy manipulált audiovizuális tartalmak képesek kihasználni az emberi agy vizuális információkra való fogékonyságát. Hameleers és munkatársai kutatása szerint a téves információk sokkal meggyőzőbbnek tűnhetnek, ha audiovizuális formában jelennek meg, szemben a szöveges tartalommal. Ez a jelenség abból adódik, hogy az audiovizuális tartalmak – amelyek egyidejűleg használják a látványt és a hangot – erőteljesebb érzelmi hatást képesek kiváltani, így nagyobb befolyást gyakorolhatnak a nézők vagy hallgatók véleményére és meggyőződéseire.³²³

Megtévesztő szándékkal készített deepfake-ek

A megtévesztő szándékkal készített deepfake tartalmak hatásai szerteágazók.³²⁴ Danielle Citron és szerzőtársa arra hívják fel a figyelmet, hogy az olyan (idő)érzékeny területeken bevetett félrevezető vizuális tartalmak, mint a nemzetközi kapcsolatok, a diplomácia, vagy a tőzsde kiemelt kockázatot képviselnek.³²⁵ Egy olyan hamisított videó, amelyben egy politikus provokatív vagy kompromittáló kijelentéseket tesz, súlyos nemzetközi konfliktusokat idézhet elő, hasonlóképpen a tőzsdén egy megtévesztő hamis nyilatkozat azonnal nemvárt árfolyammozgásokhoz vagy árfolyameséshez vezethet.

³²² Goodfellow, Ian, Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014): Generative Adversarial Nets. 2-8.o. arXiv:1406.2661

³²³ Hameleers, M, T. E. Powell, T. G. Van Der Meer, L. Bos. (2020): A Picture Paints a Thousand Lies? The Effects and Mechanisms of Multimodal Disinformation and Rebuttals Disseminated via Social Media. *Political Communication*, 37(2):281–301.o.,

³²⁴ Indiai esettanulmány XXX

³²⁵ A deepfake technológia alkalmazása zavaró hatással lehet a nemzetközi kapcsolatokra, képes manipulálni az államok közötti viszonyokat és szítani a feszültséget. Hamis videók vagy hangfelvételek készítésével, melyek diplomáciai beszélgetéseket vagy cselekedeteket ábrázolnak, a deepfake-ek fokozhatják a feszültségeket, provokálhatják a konfliktusokat és alááshatják a diplomáciai erőfeszítéseket, melynek súlyos következményei lehetnek a nemzetbiztonságra és a globális stabilitásra nézve. Például egy jól időzített deepfake, amelyet úgy terveztek, hogy a kellő időben korbácsolja fel a közvéleményt. (pl csúcstalálkozó alatt terjed el széles körben, átmenetileg tarthatatlanná téve valamelyik fél számára az egyeztetések folytatását, egy adott megállapodás elérését/bejelentését) ahogyan azt egyébként tette volna. Lásd: Chesney, Bobby, Danielle Citron (2019): Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, Vol. 107, 1775-1786.o.

A szintetikus kép és videótartalmak elterjedt változatát képviselik azon politikai deepfakek, melyeknek a közvélemény manipulálásán túl, a politikai polarizáció növelése vagy konkrét választások befolyásolása is célja lehet.³²⁶

Az orosz-ukrán háború kirobbanását követően, az orosz erők aktívan használták a deepfake technológiát abból a célból, hogy megtörjék az ukrán nép egységét és háborús elszántságát, illetve, hogy megrendítsék bizalmat a katonák, civilek és a nemzetközi közvélemény körében, ezáltal is erősítve Moszkva tárgyalási pozícióit. A legjelentősebb ilyen próbálkozás 2022 márciusában történt, amikor egy olyan hamis videó terjedt el az interneten, amelyben az ukrán elnök, Volodimir Zelenszkij felszólította katonáit, hogy tegyék le a fegyvert. A videót széles körben megosztották, számos platformon, köztük ukrán honlapokon, a Facebookon, a Twitteren, a YouTube-on és több orosz csatornán is megjelent, de a vizsgálatok hamar megerősítették, hogy MI által generált deepfake volt.³²⁷ 2023 novemberében egy újabb hasonló hamis videó látott napvilágot, ezúttal az ukrán fegyveres erők főparancsnokát, Valerij Zaluzsnijt ábrázoló manipulált felvételen, amelyben Zelenszkijhez hasonlóan ő is a fegyverletételre szólított fel. Más hamis videókban Zelenszkijt szvasztika viselete közben, illetve egy LMBTQ-felvonuláson lépve mutatták be. Noha a hivatalos orosz álláspont elutasítja a videók orosz eredetét, joggal feltételezhető, hogy ezek az Oroszország által kezdeményezett és terjesztett deepfake-ek a kiterjedt információs hadviselés és befolyásolás eszközeként szolgáltak.³²⁸

³²⁶ Az audio deepfake tartalmak előállításának közvetlen politikai kockázatai miatt döntött úgy az OpenAI hogy a saját Voice Engine nevű GenMI modelljét egyelőre nem teszi szabadon elérhetővé. A modell képes bármely 15 másodperces hangmintából a bementi forrással teljesen megegyező szintetikus hangot generálni. A hang klónozását követően a felhasználó beírhat szöveget a Voice Engine-be, és megkapja a GenMI által generált hangkimenetet (text-to-audio model). Bár rengeteg pozitív felhasználási módja lehet a jövőben -például a beszédkárosodással küzdők számára is lehetőséget nyújthat a saját hangjuk visszanyerésére - az OpenAI úgy döntött, hogy a hangklónozásból eredő potenciális veszélyekre tekintettel - kiváltképp a 2024-es globális választások évében - nem teszi még elérhetővé a modellt. Hivatalos oldalukon azt írják: Tudatában vagyunk annak, hogy az emberi beszédet imitáló tartalmak előállításának komoly kockázatai vannak, kiváltképp a választási évben. A Voice Engine-t ma tesztelő partnerek elfogadták a felhasználási irányelveinket, amelyek tiltják létező személyek és szervezetek személyazonosságának engedély nélküli vagy jogosulatlan utánzását. (We recognize that generating speech that resembles people's voices has serious risks, which are especially top of mind in an election year. The partners testing Voice Engine today have agreed to our usage policies, which prohibit the impersonation of another individual or organization without consent or legal right.) Edwards, Benj (2024): OpenAI Can Re-Create Human Voices—but Won't Release the Tech Yet. WIRED. Elérhető: <https://www.wired.com/story/openai-voice-engine-artificial-intelligence-release> (2024. 04. 03.); OpenAI (2024): Navigating the Challenges and Opportunities of Synthetic Voices. Blog, Product Announcements. Elérhető: <https://openai.com/blog/navigating-the-challenges-and-opportunities-of-synthetic-voices> (2024. 04. 02.)

³²⁷ Allyn, Bobby (2022): Deepfake Video of Zelenskyy Could Be 'Tip of the Iceberg' in Info War, Experts Warn. NPR. Elérhető: <https://www.npr.org/2022/03/16/1087062648/deepfake-video-zelenskyy-expertswar-manipulation-ukraine-russia> (2023. 06. 14.)

³²⁸ Byman, Daniel, Linna Jr., D. W., Subrahmanian, V. S. (2024): Government Use of Deepfakes. Center for Strategic & International Studies. 5-23.o.

A 2023-es szlovák parlamenti választások alatt több deepfake hangfelvétel is elterjedt a közösségi médiában, ezekben a Progresszív Szlovákia (Progresívne Slovensko) EU-párti, liberális párt elnöke, Michal Šimečka és Monika Tódová nevű újságíró között meghamisított párbeszéd volt hallható. A beszélgetés során többek között olyan témákat érintettek, mint a választási eredmények meghamisításának terve vagy a sör árának felemelése. A hanganyagokat a szavazás kezdetét megelőző 48 órás kampánycsend időszaka alatt tették közzé, emiatt - a szlovák választási szabályok értelmében - a poszt cáfolása és az azzal kapcsolatos politikai kommunikáció is akadályozott volt. Nem véletlen a célzott támadás a párt és a politikusok ellen, hiszen az akkori a legfrissebb közvélemény-kutatások szerint a Progresszív Szlovákia felzárkóztak az addigi vezető párthoz, a Smerhez. A választásnak a nemzetközi politika jelentősége is volt, mivel a Robert Fico által vezetett SMER a választási kampány során nyíltan EU ellenes retorikát folytatott, illetve megválasztása esetén Ukrajna katonai segítségük leállítását is ígérte. A videókat végül eltávolították a YouTube-ról, de a Facebookon továbbra is elérhetők maradtak. ³²⁹ (A Meta szabályzata jelenleg nem tiltja a megtévesztő információt közvetítő audio, tehát kizárólag hangalapú deepfake-eket) ³³⁰

2024 januárjában a New Hampshire-i előválasztások előtt deepfake Biden-robothívásokkal, igyekeztek eltántorítani választókat szavazatuk leadásától. A robothívás - Biden elnök hangját imitálva - arra szólította fel a szavazókat, hogy ne szavazzanak, és azt a megtévesztő üzenet hangsúlyozta, hogy a szavazatukat a novemberi választásokra kell tartalékolniuk. ³³¹

2023. május 22-én reggel rövid időre elterjedt a közösségi médiában egy hamis, GenMI által generált kép az amerikai Pentagon előtti robbanásról. Ugyan a védelmi minisztérium hivatalos

³²⁹ Meaker, Morgan (2023): Slovakia's Election Deepfakes Show AI Is a Danger to Democracy. WIRED. Elérhető: <https://www.wired.com/story/slovakias-election-deepfakes-show-ai-is-a-danger-to-democracy/?redirectURL=https%3A%2F%2Fwww.wired.com%2Fstory%2Fslovakias-election-deepfakes-show-ai-is-a-danger-to-democracy%2F> (2023. 11. 05.)

³³⁰ A vállalat várható jövőbeli tartalommoderálási politikája szerint az MI által generált videókat, képeket és hanganyagokat, címkézni fogják, hogy a felhasználók tudják, hogy ezek mesterségesen lettek létrehozva vagy módosítva. Az oversight board, amely Meta tartalompolitikai felülvizsgálatait végzi, javasolta, hogy a cég bővítse a közvélemény megtévesztésére szánt video deepfakek tiltását az audio deepfake-ekre is. Hozzáteszik, hogy a manipulált tartalmak eltávolítását csak a legmagasabb kockázatú esetekre kell korlátozni, ahol a tartalom egyértelműen kárt okoz. Frissítés: A Meta 2024 április 5-én jelentette be, hogy jelentős változtatásokat fog eszközölni a digitálisan létrehozott és módosított tartalmakra vonatkozó szabályzatában. A vállalat májustól Made by AI (Mesterséges intelligenciával készült) címkéket fog alkalmazni a Facebookon és az Instagramon közzétett, GenMI-vel generált videókra, képekre és hanganyagokra is, lényegesen kiterjesztve korábbi szabályzatát mely csak a manipulált videók egy szűk szeletére vonatkozott. Lásd: Meta, Bickert, Monika (2024): Our Approach to Labeling AI-Generated Content and Manipulated Media. META Newsroom. Elérhető: [https://about.fb.com/news/2024/04/metas-approach-to-labeling-ai-generated-content-and-manipulated-media/\(2024. 04. 05.\)](https://about.fb.com/news/2024/04/metas-approach-to-labeling-ai-generated-content-and-manipulated-media/(2024. 04. 05.))

³³¹ Eliott, Vittoria (2024): The Biden Deepfake Robocall is Only the Beggining. Wired. Elérhető: <https://www.wired.com/story/biden-robocall-deepfake-danger/> (2024. 01. 31.)

közleményében hamar cáfolta a kép valóságtartalmát, a robbantásról szóló hamis hírek rövid időre megingatták a főbb részvénypiaci indexeket. Külön nehezítette a helyzetet, hogy a híresztelést terjesztő Twitter-fiókok többségén kék pipajel volt, ami korábban a hitelesített fiókokat jelezte, de Elon Musk felvásárlását követő változtatások nyomán mára bárki megvásárolhatja a Twitter Blue előfizetéssel. A hamis Pentagon-képet megosztó fiókok között volt egy a Bloomberg hírügynökséget utánozó fiók is, valamint a valódi, Kreml-közeli orosz RT hírszolgálat fiókja. Az RT később törölte a bejegyzését, míg a hamis Bloomberg-fiókot felfüggesztette a Twitter.³³²

Explicit deepfake tartalmak

Egy ritkábban tárgyalt, de szintén aggodalomra okot adó trend a GenMI eszközök segítségével létrehozott szexuális tartalmú deepfake képek és videók rohamos terjedéséhez köthető. Ilyen tartalmak előállítása bármiféle grafikus tudás vagy videószerkesztői előismeret nélkül lehetséges a lényegében bárki által hozzáférhető és nyílt forráskódú technológiák (pl. a CivitAI szoftver) segítségével.³³³ A GenMI eszközök felbukkanásával drámai növekedés volt megfigyelhető az explicit deepfake videók számában. Míg 2022-ben az interneten körülbelül 3725 ilyen tartalom volt elérhető, ez a szám 2023 21 019-re emelkedett. A New York Times egy 2024-es cikkében egyenesen járványnak nevezte a jelenséget.³³⁴

Akárcsak a politikai deepfakek esetében, az explicit tartalmak áldozatául is eshetnek közéleti személyek és politikusok. Giorgia Meloni, Olaszország miniszterelnöke 100 000 euró kártérítést követelt, miután deepfake pornográf videókat készítettek és tettek közzé róla az interneten a hozzájárulása nélkül. A videókban Meloni arcát egy másik személy testére illesztették digitális módszerekkel, majd feltöltötték őket egy amerikai weboldalra.³³⁵

³³² Brian Bushard (2023): Fake Image Of Explosion Near Pentagon Went Viral—Even Though It Never Happened. Forbes. Elérhető: <https://www.forbes.com/sites/brianbushard/2023/05/22/fake-image-of-explosion-near-pentagon-went-viral-even-though-it-never-happened/>

³³³ A CivitAI platformján több ezer szintetikus színész/színésznő érhető el (ezek gyakran valós emberek képeiből készülnek, amelyeket sok esetben engedély nélkül gyűjtenek össze az internet különböző részeiről.) és bár a platform ezt nyilvánosan nem reklámozza, amennyiben a szolgáltatást igénybe vevő ilyen célra kívánja azt használni, pornográf tartalmak előállítása is könnyen lehetségessé válik. Sőt a felhasználó a saját eszközén tárolt, általa ismert személyről készült képet is felhasználhatja hasonló célra. A témáról, és a CivitAI platformján elérhető erősen szexuális tartalmú beágyazott szerkesztőfilterekről bővebben: Maiberg Emanuel (2023): Inside the AI Porn Marketplace Where Everything and Everyone Is for Sale. 404media. Elérhető: <https://www.404media.co/inside-the-ai-porn-marketplace-where-everything-and-everyone-is-for-sale/>

³³⁴ A szexuális tartalmú deepfake-k túlnyomórészt nőket céloznak meg, egy felmérés szerint az ilyen videókban szereplő alanyok 99%-a nő.

³³⁵ A hatóságok videók elkészítésével egy 40 éves férfit és 73 éves apját gyanúsítják, akiket mobileszköz-nyomkövetéssel azonosítottak Meloni ügyvédje, Maria Giulia Marongiu azt nyilatkozta, hogy amennyiben a bíróság megítéli ügyfele számára a 100 000 eurós kártérítési összeget a miniszterelnök olyan nők támogatására

A megtévesztésen túl

Egy az Egyesült Királyságban végzett nagymintás felmérés (N=2005), amely a hamisított digitális tartalmakkal kapcsolatos állampolgári attitűdöket vizsgálta, kimutatta, hogy a deepfake technológiák jó eséllyel növelik az általános bizonytalanságot, szkepticizmust és cinizmust a lakosság körében. A felmérés megállapította, hogy az internethasználók egy része hiába volt képes felismerni a deepfake videókat, ezeknek a hamisítványoknak a pusztja jelenléte számottevően csökkentette az online hírekbe vetett általános bizalmukat. A tanulmány írói hangsúlyozták, hogy az elbizonytalanodás érzése nem korlátozódik csupán az adott deepfake videót tartalmazó híradásra, hanem általánosságban véve is csökkenti az online megjelenő hírekbe vetett bizalmat (kiváltképp a közösségi médiában megjelenő híreket)³³⁶

A deepfake-ek megjelenése valójában csak tovább fokoz egy olyan trendet, melynek a jelei már korábban is megjelentek. A híradásokba vetett bizalom az elmúlt egy évben további 2 százalékponttal csökkent. A teljes mintánkból tíz emberből átlagosan csak négyen (40%) mondják azt, hogy a legtöbbször megbíznak a hírekben.³³⁷ E tekintetben is az idősebb generáció tagjai a legkiszolgáltatottabbak.³³⁸

A 2023-as évben tapasztalt innovációshullám bizonyítékul szolgált arra, hogy a jövőben a jelenleginél még inkább fejlettebb és mindenki számára könnyen hozzáférhető modellek, szolgáltatások jelennek majd meg az online térben. Ahogy a GenMI technológiák fejlődnek, a deepfake-ek egyre valóságghűbbek és nehezebben észlelhetőek lesznek. Ez növeli az előbb ismertetett politikai és társadalmi kockázatokat, valamint szükségessé teszi mind a kifinomult felismerési és szűrési technológiák fejlesztését, mind pedig a kockázatokat adekvát módon kezelni képest adó jogi keretrendszerek megalkotását.

kívánja fordítani, akik férfiak általi erőszak áldozatai lettek. Ez az összeg szimbolikus jelentőségű, és célja, hogy üzenetet küldjön a hasonló visszaélések áldozatainak, ne féljenek feljelentést tenni Gozzi, Laura (2024): Giorgia Meloni: Italian PM seeks damages over deepfake porn videos. Elérhető: <https://www.bbc.com/news/world-europe-68615474> (2024. 03. 21.); Starcevic, Seb (2024): Italy's Giorgia Meloni called to testify in deepfake porn case. Elérhető: <https://www.politico.eu/article/italian-pm-giorgia-meloni-called-to-testify-in-deepfake-porn-case/> (2024. 03. 21)

³³⁶ Vaccari, C., & Chadwick, A. (2020). Deepfakes and Disinformation: Exploring the Impact of Synthetic Political Video on Deception, Uncertainty, and Trust in News. *Social Media + Society*, 6(1) <https://doi.org/10.1177/2056305120903408>.

³³⁷ Reuters Institute (2023): Digital News Report. 9-10.o. Elérhető: <https://reutersinstitute.politics.ox.ac.uk/digital-news-report/2023>

³³⁸ Egy 2019-ben végzett felmérés azt mutatta, hogy a Facebook (Meta) közösségi média felületét használó idősebb korosztály tagjai hajlamosabbak arra, hogy elhiggyék és megosszák a téves információkat. Ennek legvalószínűbb oka, hogy a most 60-as éveikben járó és azon túli korosztály nem rendelkezik a digitális médiaműveltségnek azon szintjével, amely ahhoz szükséges, hogy meg tudja állapítani az online hírek valóságtartalmát/ megbízhatóságát.

6.10 Az Európai Unió dezinformációs stratégiája

Az Európai Unió az elmúlt években jelentős erőfeszítéseket tett a dezinformáció elleni küzdelem terén és számos olyan intézkedést hozott, amelyek célja a digitális tér biztonságának növelése és a demokratikus folyamatok védelme volt. Az Unió ezirányú törekvései 2015 márciusáig vezethetők vissza, ekkor bízta meg az Európai Tanács az Unió külügyi és biztonságpolitikai főképviselőjét, hogy dolgozzon ki egy *stratégiai kommunikációra vonatkozó cselekvési tervet* Oroszország folyamatos dezinformációs kampányainak kezelésé. E cselekvési terv eredményeként hoztak létre egy új stratégiai kommunikációs részleget (StratCom) az Európai Külügyi Szolgálaton belül. A StratCom első munkacsoportjának - az úgynevezett keleti stratégiai kommunikációs munkacsoportnak - feladatává vált az Unión kívülről (elsősorban Oroszországból) származó dezinformáció elleni fellépés, valamint pozitív stratégiai kommunikáció kialakítása és terjesztése az Unió keleti szomszédságában.³³⁹

A Bizottság 2017 végén hozta létre az Álhírekkel és dezinformációval foglalkozó magas szintű szakértői csoportot (High-Level Expert Group on Fake News and Disinformation), melynek legfőbb feladatául szabták, hogy konkrét, a gyakorlatban alkalmazható tanácsokkal szolgáljon a dezinformáció kezeléséhez. A csoport a rákvetkező év márciusában nyújtotta be jelentését, amelyben többek között az átláthatóság növelésével, a felhasználók és hírolvasók médiatudatosság elősegítésével, illetve a platformszolgáltatókkal közös hatékony együttműködéssel kapcsolatban fogalmaztak meg javaslatokat.³⁴⁰

A jelentés fontos alapjául szolgált a Bizottság által ugyanazon év áprilisában közzétett „megközelítés az online félretájékoztatás kezelésére” című közleményének.³⁴¹ E

³³⁹ A Kelet-StratCom munkacsoport működteti például az EUvsDisinfo weboldalt is, melynek elsőszámú célja, hogy felfedje az orosz forrásból származó hamis állításokat. A megfogalmazott hosszútávú célok között szerepel ezenfelül, a Kreml információs műveleteivel kapcsolatos állampolgári tudatosság növelése, valamint a polgárok megfelelő tudással történő felvértezése a megtévesztő digitális információkkal és a média manipulációjával szemben. A weboldal az alábbi linken érhető el: <https://euvsdisinfo.eu/>

³⁴⁰ A javaslat az átláthatóság növelése érdekében szorgalmazta például a tájékoztatásnyújtási kötelezettség kiterjesztését a digitális médiaplatformokon megjelenő hirdetések forrásai és finanszírozási módjaival kapcsolatban. A médiatudatosság növelése érdekében a javaslat az újságírói függetlenség biztosításán, valamint a tényellenőrző szervek támogatásán túl különböző oktatási programok indítását is javasolta, melyek arra tanítanak a hírfogyasztó állampolgárokat, hogyan értékeljék kritikusan az információkat, hogyan ismerjék fel a dezinformációt, és hogyan navigáljanak sikeresen az újonnan létrejövő online környezetben. A digitális platformokkal való együttműködéshez kötődően olyan intézkedések közös kidolgozását és folyamatos monitorozásával javasolták, melyek nagyobb felelősségvállalásra köteleznék online platformokat a dezinformációs tartalmak terjedésének csökkentésében. Erről bővebben: European Commission 2018: A multi-dimensional approach to disinformation. Report of the independent High level Group on fake news and online disinformation. Luxembourg, Publications Office of the European Union, 22-31o.

³⁴¹ Európai Bizottság (2018): Európai megközelítés az online félretájékoztatás kezelésére. COM(2018) 236 final

megközelítésben a Bizottság megállapítja, hogy a *szigorú szakmai előírásokon alapuló, jól működő, szabad és plurális információs környezet* elengedhetetlen feltétele az egészséges demokratikus viták létrejöttének, egyúttal pedig a demokratikusan működő társadalmak hosszútávú fenntartásának. Ennek előmozdítása érdekében ajánlásokat fogalmazott meg a *magas színvonalú információk előnyben részesítésén alapuló* digitális ökoszisztémák megteremtésére.³⁴² A szöveg ezenfelül szorgalmazza egy európai gyakorlati kódex kidolgozását is a dezinformáció visszaszorítása céljából.

Az uniós erőfeszítések kiemelt jelentőséget szenteltek a dezinformáció visszaszorítását célzó, önszabályozó gyakorlati kódex megalkotására. Az EU és a nagy online platformok között létrejött önkéntes megállapodás legfőbb célkitűzése, hogy átláthatóbbá és elszámoltathatóbb tegye a legnagyobb befolyással bíró online platformokat, illetve, hogy egyfajta innovatív keretként szolgáljon ezen platformok dezinformációval kapcsolatos politikáinak nyomon követéséhez és javításához.³⁴³

Az Európai Tanács a 2018. júniusi következtetéseiben kérte fel a Bizottságot arra, hogy *terjesszenek elő konkrét javaslatokat a félretájékoztatás jelentette problémára adandó koordinált uniós válaszingedményekre vonatkozóan*, valamint határozzák az e célból biztosítandó szükséges mértékű és elégséges erőforrásokat.³⁴⁴

Az Európai Bizottság 2018 decemberében egy cselekvési tervet is közzétett a dezinformáció elleni küzdelem elősegítésére, melyben az Unió által megfogalmazott ajánlások az alábbi négy fő célra összpontosítottak: (1) Az uniós intézmények azon képességeinek javítása, hogy felderítsék, elemezzék és leleplezzék a dezinformációt (2) A dezinformációval kapcsolatos koordinált és együttes válaszlépések erősítése (3) A magánszektor bevonása és mozgósítása a dezinformáció elleni küzdelemben (4) A dezinformációból eredő problémák tudatosítása és a társadalom ellenálló képességének javítása.³⁴⁵

³⁴²A megfogalmazott ajánlások az alábbi öt területbe sorolhatóak: (1) Átláthatóbb, megbízhatóbb és elszámoltathatóbb online környezet kialakítása (2) Biztonságos és ellenállóképes választási rendszerek és folyamatok garانتálása (3) Az oktatás és a médiaműveltség előmozdítása (4) A színvonalas és minőségi újságírás támogatása (5) A belső és külső információs fenyegetések elleni kommunikációs stratégiák kidolgozása. Lásd bővebben: Európai Bizottság (2018): Európai megközelítés az online félretájékoztatás kezelésére. COM(2018) 236 final 7-18o.

³⁴³ Az önszabályozási megközelítésen alapuló kódexben a 16 aláíró fél összesen 21 kötelezettséget vállalt többek között olyan intézkedésekkel kapcsolatban mint; reklámelhelyezések ellenőrzése, politikai és a közügyekkel kapcsolatos hirdetések átláthatósága, szolgáltatások integritása, tudatos fogyasztói magatartás előmozdítása, tényellenőrzők és a kutatók szerepvállalásának erősítése, kódex eredményességének mérése.

³⁴⁴ Európai Tanács 2018: Következtetések. EUCO 9/18 (2018. június 28.) II. Cím (Biztonság és védelem), 10. bekezdés, 5. oldal

³⁴⁵ European Commission (2018): Action Plan against Disinformation. JOIN(2018) 36 final, 5-11.o.

6.11 Digitális Szolgáltatásokról Szóló Rendelet

Az Európai Unió új Digitális Szolgáltatásokról szóló rendelete (DSA) ³⁴⁶ - amely az elektronikus kereskedelemről szóló irányelvet váltotta fel - harmonizált jogszabályokat állapít meg az Unió állampolgárainak életére közvetlen hatással bíró³⁴⁷ online közvetítő szolgáltatók számára. A rendelet célja egy biztonságos, kiszámítható és megbízható online környezetet megteremtése, amely többek között hatékonyan képes kezelni a jogellenes online tartalmak, a személyre szabott információs terek és a dezinformációs kampányok terjedéséből eredő társadalmi kockázatokat. ³⁴⁸

Az Európai Bizottság 2020 decemberében terjesztette elő a DSA tervezetét egy átfogó konzultációs és hatásvizsgálati folyamat után. A javaslatot ezután az Európai Parlament és az EU Tanácsa tárgyalta meg és módosította a rendes jogalkotási eljárás keretében. Hosszas egyeztetések után 2022. október 19-én hivatalosan is elfogadták és kihirdették a végleges szöveget mely ezáltal formálisan is az EU jogrendjének részévé vált. A szabályozás nagy része, beleértve az összes online platformra vonatkozó kötelezettségeket, 2024. február 17-től vált kötelezően alkalmazandóvá, 15 hónapos felkészülési időt biztosítva a vállalkozásoknak, hogy hozzáigazítsák működésüket az új szabályokhoz. A nagyon nagy online platformokra (Very Large Online Platforms, VLOPs) és a nagyon nagy online keresőprogramokra (Very Large Online Search Engine, VLOSEs) vonatkozó szigorúbb kötelezettségek már 2023. augusztus 25-től érvénybe léptek.

A DSA egy többszintű és differenciált szabályozási modellt vezetett be, amely annak alapján különbözteti meg az egyes szolgáltatókat, hogy azok mekkora felhasználói bázissal rendelkeznek. Ez a modell tükrözi azt a felismerést, hogy a nagyobb platformok jelentősebb hatást gyakorolnak az alapvető jogok érvényesülésére, a politikai diskurzusra és a demokratikus

³⁴⁶ Európai Bizottság: Az Európai Parlament és a Tanács rendelete a digitális szolgáltatások egységes piacáról (digitális szolgáltatásokról szóló jogszabály) és a 2000/ 31/EK irányelv módosításáról. COM(2020)825 final

³⁴⁷ A DSA követelményeinek alkalmazását az igénybe vevők lakó- vagy tartózkodási helyéhez köti, nem pedig a közvetítő szolgáltató letelepedési helyéhez vagy székhelyéhez. A 2. cikk (1) bekezdése kimondja, hogy *a rendelet az olyan igénybe vevők részére kínált közvetítő szolgáltatásokra vonatkozik, amelyeknek, illetve akiknek a letelepedési helye az Unióban található vagy akik, illetve amelyek az Unióban találhatók, függetlenül az említett közvetítő szolgáltatást biztosító szolgáltatók letelepedési helyétől.*

³⁴⁸ DSA Preambulum 9.

nyilvánosságra. Ebből következik, hogy nagyobb a felelősségük is a társadalmi hatások kezelése tekintetében. Ennek értelmében a rendelet szigorúbb jogi követelményeket állapít meg az ún. online óriásplatformok és nagyon népszerű online keresőprogramok számára. Azon szolgáltató tartoznak e kettő kategóriába, amelyek havonta átlagosan legalább 45 millió aktív igénybe vevővel rendelkeznek az EU-ban (jelentősebb és gyorsan bekövetkező népességváltozás esetén az Unió összlakosságának 10% a mértékadó) ³⁴⁹

Ezekre a platformokra vonatkozó további kötelezettségek közé tartozik például az általuk nyújtott szolgáltatások használata során felmerülő online veszélyekkel kapcsolatos kockázatok átfogó éves értékelése. Az ilyen platformoknak biztosítaniuk kell, hogy rendszeresen felülvizsgálják és mérsékeljék a felhasználókat érintő potenciális kockázatokat, beleértve a felhasználói jogokra, az adatvédelemre és a szólásszabadságra gyakorolt hatásokat. ³⁵⁰

A dolgozat témájához kötődően a DSA jogszabály következő rendelkezéseit indokolt ismertetni:

6.11.1 Átláthatósági és tájékoztatási kötelezettség

Az alapvető jogok védelme tekintetében kiemelkedő jelentőségű 14. cikk, amely amellett, hogy szigorú transzparenciakövetelményeket állít, széles körű tájékoztatási kötelezettséget is előír a közvetítő szolgáltatók számára. Ennek alapján *világosan, egyszerűen, érthetően, felhasználóbarát módon* tájékoztatást kell nyújtani a szolgáltatást igénybe vevők számára,

³⁴⁹ A Bizottság 2023. április 25-én azonosította azon szolgáltatókat, amelyek e két kategória valamelyikébe tartoznak. Ennek értelmében óriásplatformnak (Very Large Online Platform, VLOP) minősül az: AliExpress / Amazon Store / AppStore / Booking / Facebook / Google Maps / Google Play / Google Shopping / Instagram / LinkedIn / Pinterest / Snapchat / TikTok / Twitter / Wikipedia / YouTube / Zalando. Nagyon népszerű keresőmotor szolgáltatóként (Very Large Searching Engine, VLSE) pedig a Bing és Google Search szolgáltatókat nevezték meg..

³⁵⁰ A szolgáltatóknak a rendszerszintű kockázatok négy meghatározott kategóriáját kell részletesen értékelniük: (1) jogellenes tartalmak terjesztésével (pl.: gyűlöletbeszéd) vagy a jogellenes tevékenységek folytatásával kapcsolatos kockázatok; (2) a felhasználók alapvető jogainak gyakorlását érintő tényleges vagy várható negatív hatások (pl.: véleménynyilvánítás és a tájékozódás szabadságát, tömegtájékoztatás szabadságát vagy a magánélethez való jogot érintő korlátozások); (3) a demokratikus folyamatokra, a polgári közbeszédre, a választási folyamatokra és a közbiztonságra gyakorolt tényleges vagy előrelátható negatív hatások; (4) a nemi alapú erőszakkal, a közegészség és a kiskorúak védelmével összefüggő tényleges vagy előre látható negatív hatások, valamint a személyek testi és szellemi jóllétére gyakorolt súlyos negatív következmények (pl.: manipuláción keresztül vagy olyan online felületek kialakítása útján, melyek függőségéhez vezethetnek) DSA 34. cikk, Preambulum 80-84. DSA 33-35. cikk

többek között a tartalommoderálás³⁵¹ céljából alkalmazott valamennyi szabályra, eljárásra, intézkedésre és eszközre – ideértve az algoritmikus döntéshozatalt és az emberi felülvizsgálatot is.³⁵²

6.11.2 A személyre szabott célzott hirdetés és manipuláció tilalma

A DSA szintén fontos eleme, hogy az online platformot üzemeltető szolgáltatók nem tervezhetik meg, alakíthatják ki vagy üzemeltethetik online interfészeiket oly módon, *amely megteveszti vagy manipulálja a szolgáltatásaikat igénybe vevőket, vagy gyengíti azok szabad és tájékozott döntéshozatalra való képességét.*³⁵³

A 26. cikk (3) bekezdése kimondja, hogy az online platformokon nem jeleníthetnek meg profilalkotáson alapuló hirdetéseket a szolgáltatások igénybe vevői számára a személyes adatoknak a GDPR-ban említett különleges kategóriáinak felhasználásával. Ide tartoznak a faji vagy etnikai származásra, politikai véleményre, vallási vagy világnézeti meggyőződésre vagy szakszervezeti tagságra utaló személyes adatok, valamint a természetes személyek egyedi azonosítását célzó genetikai és biometrikus adatok, az egészségügyi adatok és a természetes személyek szexuális életére vagy szexuális irányultságára vonatkozó személyes adatok.³⁵⁴

A DSA 26. cikke értelmében az online hirdetéseket megjelenítő platformok üzemeltetőinek *világosan, tömören, egyértelműen és valós időben fel kell tüntetniük, amennyiben az általuk közölt információ hirdetésnek minősül.* Emellett fel kell tüntetniük azt is, hogy mely természetes vagy jogi személy nevében került sor a hirdetés megjelenítésére 26. cikk 1(b); a hirdetést mely természetes vagy jogi személy finanszírozta, amennyiben e személy eltér a b) pontban említett

³⁵¹ DSA 3.cikk t) értelmében a tartalommoderálása a közvetítő szolgáltató olyan automatizált vagy nem automatizált tevékenysége, amely különösen a szolgáltatás igénybe vevője által közzétett jogellenes tartalom vagy a közvetítő szolgáltató szerződési feltételeivel összeegyeztethetetlen információ észlelésére, azonosítására és kezelésére szolgál, ideértve az ilyen jogellenes tartalom vagy információ elérhetőségét, láthatóságát és hozzáférhetőségét érintő intézkedéseket, például annak hátrасorolását, a demonetizálását, az ahhoz való hozzáférés megszüntetését vagy annak eltávolítását, vagy a szolgáltatás igénybe vevője általi információközlés lehetőségét érintő intézkedéseket, például a fiókja megszüntetését vagy felfüggesztését

³⁵² Ezen felül, a DSA 14. cikke azt is kimondja, hogy a szolgáltató az ÁSZF jelentősebb módosításairól külön köteles tájékoztatni az igénybe vevőket, ha pedig a szolgáltatás elsődlegesen vagy túlnyomórészt a kiskorúakat célozza, akkor számukra érthető nyelvezettel is ismertetni kell az igénybevételi feltételeket. Az online óriásplatformot és a nagyon népszerű online keresőprogramot üzemeltető szolgáltatóknak külön kötelezettségük, hogy a sokszor igen összetett felhasználási feltételekről összefoglalót készítsenek és tegyenek közzé könnyen hozzáférhető és géppel olvasható formátumban, és azt az összes olyan tagállam nyelvén meg is jelenítsék, ahol szolgáltatást nyújtanak .

³⁵³ DSA 25. cikk (1)

³⁵⁴ GDPR 9.cikk (1))

természetes vagy jogi személytől 26.cikk 1(c); valamint azt is, hogy mely paraméterek szolgálták a hirdetéssel megcélzott személyek meghatározására 26. cikk 1(d).³⁵⁵

Az online környezetben az egyének számára a legtöbb esetben nincs lehetőség megérteni, hogy hogyan és miért jut el hozzájuk bizonyos tartalom. Ezt az információs asszimetriát orvosolja a DSA azzal, hogy kötelezi az összes online platformot, hogy közzé tegye tartalomajánló rendszereik³⁵⁶ paramétereit, ezáltal világos magyarázatot szolgáltatva arra, hogy miért is ajánlják az adott információkat a szolgáltatás igénybe vevője számára.³⁵⁷

Ezen túlmenően a 38. cikk előírja, hogy az ajánlórendszereket használó online óriásplatformot vagy nagyon népszerű online keresőprogramot üzemeltető szolgáltatók minden ajánlórendszerük esetében legalább egy olyan lehetőséget biztosítanak, amely nem a GDPR szerint meghatározott profilalkotási³⁵⁸ eljáráson alapul. A DSA nem teszi kötelezővé a nem profilozáson alapuló ajánlást, de ösztönözheti a platformokat, hogy kínáljanak ilyen lehetőséget.

A rendelet olyan új mechanizmusokat is bevezet, amelyek segítenek a felhasználóknak bejelenteni a jogellenes tartalmakat³⁵⁹ annak érdekében, hogy a platformok hatékonyan azonosíthatják és eltávolíthatják azokat. A DSA hangsúlyozza, hogy a platformoknak proaktív intézkedéseket kell tenniük az illegális tartalmak ellen, de nem kötelezi őket általános

³⁵⁵ A 39. cikk továbbiakban előírja azt is, hogy az online hirdetéseket megjelenítő óriásplatformot és nagyon népszerű online keresőprogramot üzemeltető szolgáltatók nyilvánosan hozzáférhetővé teszik a hirdetésekkel kapcsolatos további információkat is. 39. cikk (2) a) a hirdetés tartalma, beleértve a termék, a szolgáltatás vagy a márka nevét és a hirdetés tárgyát; b) kinek nevében került sor a hirdetés megjelenítésére; c,) mely természetes vagy jogi személy fizette a hirdetést, amennyiben e személy eltér a b) pontban említett személytől; d) mely időszakban került sor a hirdetés megjelenítésére

³⁵⁶ DSA 3. cikk s) „ajánlórendszer”: teljes mértékben vagy részben automatizált rendszer, amelyet az online platform arra használ, hogy az online interfészén konkrét információkat javasoljon vagy ezeket az információkat rangsorolja a szolgáltatás igénybe vevője számára, többek között a szolgáltatás igénybe vevője által indított keresés alapján vagy egyéb módon meghatározva a megjelenített információk relatív sorrendjét vagy elsőbbségét.

³⁵⁷ DSA 27.cikk (1-2)

³⁵⁸ A GDPR 3. cikk (4) pontja úgy határozza meg a profilalkotást, mint a: *személyes adatok automatizált kezelésének bármely olyan formája, amelynek során a személyes adatokat valamely természetes személyhez fűződő bizonyos személyes jellemzők értékelésére, különösen a munkahelyi teljesítményhez, gazdasági helyzethez, egészségi állapothoz, személyes preferenciákhoz, érdeklődéshez, megbízhatósághoz, viselkedéshez, tartózkodási helyhez vagy mozgáshoz kapcsolódó jellemzők elemzésére vagy előrejelzésére használják*

³⁵⁹ 3. cikk h) „jogellenes tartalom”: bármely olyan információ, amely önmagában vagy egy tevékenységgel kapcsolatban, beleértve a termékek értékesítését vagy a szolgáltatások nyújtását, nem felel meg az uniós jognak vagy bármely tagállam – az uniós joggal összhangban álló – jogának, függetlenül az adott jog pontos tárgyától vagy jellegétől

monitorozásra.³⁶⁰ Ehelyett a platformoknak együtt kell működniük a "trusted flaggers" (megbízható bejelentők) rendszerével, hogy hatékonyan kezeljék az illegális tartalmakat

6.12 Észrevételek

XXX--XXX

A tartalomajánló rendszerekkel kapcsolatban várható még vita az Európai Unióban. Ezt jól szemlélteti, hogy 2023. december 20-án 17 európai parlamenti képviselő levelet írt a Bizottságnak, melyben szorgalmazták, hogy a technológiai platformok perszonalizált tartalomajánlórendszereit alapértelmezésben kapcsolják ki. Levelükben úgy fogalmaznak: *Az interakción alapuló ajánlórendszerek, különösen a kifinomult személyre szabott rendszerek komoly veszélyt jelentenek polgárainkra és társadalmunkra általában, mivel az érzelmekre ható és szélsőséges tartalmakat helyezik előtérbe, és kifejezetten a provokációra hajlamos személyeket célozzák meg.*³⁶¹ Annak ellenére, hogy ez az ötlet már a DSA tárgyalások során is felmerült, végül nem került be a végleges rendeletbe. Ehelyett a jogalkotók arra az előbb ismertetett kompromisszumos megoldásra jutottak, hogy a nagyobb platformoknak minden ajánlórendszerük esetében legalább egy olyan tartalomszolgáltatást kell nyújtaniuk, amely nem profilalkotáson alapul (opt out megközelítés).

A rendelet hiányosságaként felróható Hacker és társai azon kritikája, miszerint a szabályozás érdemi része nem terjed ki a „privát üzenetküldő szolgáltatásokra”.³⁶² Ennek következtében az olyan népszerű közösségi platformok, mint a WhatsApp és Telegram zárt csoportjai, ahol a problémás tartalmak és a dezinformáció különösen elterjedt, nem esnek a DSA szabályozási hatálya alá.³⁶³

Előfordulhat az is, hogy a jogszabályi kötelezettségek betartásának nehézségeit nem az érintett szolgáltatók jogkövetési hajlandóságának vagy a részletes, következetes szabályozás

³⁶⁰ Az általános monitorozás azt jelentené, hogy a platformok folyamatosan figyelnék minden felhasználói tartalmat, hogy illegális vagy káros tartalmakat találjanak. Ez a követelmény túlzott terhet jelentene a platformok számára, és aggályokat vetne fel a felhasználói magánélet és szólásszabadság védelme szempontjából.

³⁶¹ "Interaction-based recommender systems, in particular hyper-personalised systems, pose a severe threat to our citizens and our society at large as they prioritize emotive and extreme content, specifically targeting individuals likely to be provoked." Lásd bővebben: Natasha Lomas (2023): Big Tech's divisive 'personalization' attracts fresh call for profiling-based content feeds to be off by default in EU. TechCrunch. Elérhető: <https://techcrunch.com/2023/12/20/dsa-recommender-systems/> (2023. 12. 20.)

³⁶² DSA 14. preambulumbekzdés

³⁶³ A DSA mechanizmusainak teljes skálája csak a hagyományos közösségi hálózatokon közzétett GenMI tartalomra alkalmazható. Lásd: Hacker et al.(2023): Regulating ChatGPT and other Large Generative AI Models. Fairness, Accountability, and Transparency (FAccT '23). 1112-1123.o. <https://doi.org/10.1145/3593013.3594067>

meglétének hiánya, hanem sokkal inkább a rendelkezésre álló technológiai képesség hiánya okozza. A Nagy Nyelvi Modellek esetében például, sok esetben hiába építenek be olyan visszaélés-megelőző funkciókat és biztonsági korlátokat, melyek célja, hogy visszautasítsák vagy, megtagadják a nem kívánt utasításokat (például gyűlöletkeltő, erőszakra felbujtó esszék és pamfletek előállítása), a felhasználók mégis megtalálják a módot ezeknek a biztonsági intézkedéseknek a megkerülésére.³⁶⁴ Egy ilyen megkerülési technika például a "prompt injection" támadás, ahol úgy álcázzák a bemeneti szöveget, mint a felhasználó vagy a fejlesztő utasításait, annak érdekében felülírják a modellekbe épített vagy tanított korlátozásokat.³⁶⁵ A másik technológiai probléma a szintetikus tartalmak azonosításához köthető. Ahogy ezek a modellek egyre kifinomultabbá válnak, úgy lesz egyre nehezebb azonosítani azokat. Ennek okán a jogi szabályozás és keretrendszerek megalkotásán túl fontos a technológiai eszköztár fejlesztése és biztosítása is.³⁶⁶

Látható, hogy az Európai Unió a dezinformáció és félretájékoztatás visszaszorítása érdekében olyan szabályozási megközelítést alkalmaz, amely az online platformok működésének átláthatóságát és elszámoltathatóságát kívánja előmozdítani, (például a felhasználók tájékoztatására vonatkozó kötelezettségek garantálása révén). Ez az elképzelés tükröződik a DSA rendelkezéseiben és hasonlóképpen az MI-rendelet is ezen elv mentén lett megalkotva a GenMI által előállított szintetikus tartalmak vonatkozásában:

Az MI rendelet 50. cikk (1) bekezdése például kimondja, hogy a szolgáltatóknak biztosítaniuk kell, hogy a *természetes személyekkel való közvetlen interakcióra szánt MI-rendszereket* úgy tervezzék meg és fejlesszék ki, hogy az érintett természetes személyek tájékoztatást kapjanak

³⁶⁴ A tartalomszűrők (content filter) feladata, hogy észleljék és eltávolítsák vagy blokkolják a meghatározott kategóriák szerinti nemkívánatos tartalmakat. A biztonsági szűrők (safety filter) célja, hogy megakadályozzák az ellenséges promptokat és biztosítsák, hogy a modell kimenete biztonságos maradjon különböző támadási forgatókönyvek alatt. Ezek a szűrők összetettebb manipulációkat és ellenséges támadásokat kezelnek, amelyek megpróbálják kijátszani a modell védelmi mechanizmusait

³⁶⁵ Lásd: Fábio, Perez, Ian Ribeiro (2022): Ignore previous prompt: Attack techniques for language models. 3-6.o. arXiv:2211.09527; Az ilyen gyakorlatokhoz köthetőek az ún. jailbreaking támadások is. A kettő közötti különbség abban rejlik, hogy míg a prompt injection egy adott utasítással törekszik tiltott, vagy káros kimenet generálására, addig a jailbreaking végrehajtójának a modell biztonsági mechanizmusainak teljes felszámolása a célja. Bővebben: Yi Liu, Gelei Deng, Zhengzi Xu, Yuekang Li, Yaowen Zheng, Ying Zhang, Lida Zhao, Tianwei Zhang, Kailong Wang, and Yang Liu (2023): Jailbreaking ChatGPT via prompt engineering: An empirical study. 4-9.o. arXiv:2305.13860v2

³⁶⁶ Ide értendő többek között a digitális vízjelek és időbélyegzők, valamint a metaadat és blokklánc alapú hitelesítési technológiák fejlesztése.

arról, hogy egy MI-rendszerrel állnak interakcióban.³⁶⁷ Hasonlóan az átláthatóság és a tájékoztatási kötelezettség fontosságát tükrözi az is, hogy *a szintetikus hang-, kép-, video- vagy szöveges tartalmat létrehozó MI-rendszerek* – köztük az általános célú MI-rendszerek – szolgáltatóinak *biztosítaniuk kell, hogy az MI rendszer kimeneteit géppel olvasható formátumban jelöljék meg, és azok mesterségesen létrehozottként vagy manipuláltként észlelhetők legyenek.*³⁶⁸

Ugyanezen elv mentén írja elő a rendelet, hogy az olyan MI-rendszerek alkalmazóinak, amelyek *eredetinek vagy valóságosnak tűnő („deepfake”) kép-, hang- vagy videotartalmat hoznak létre vagy manipulálnak*, közölniük kell, hogy a tartalmat mesterségesen hozták létre vagy manipulálták. Azonban – feltehetően - annak okán, hogy a szintetikus tartalmak esetén a jogalkotók figyelme elsősorban a közvélemény manipulálásából, választások befolyásolásából, és politikai polarizáció kiéleződéséből fakadó kockázatokra összpontosul, a bekezdés második fele enyhít ezen a kötelezettségen, és azt írja: *amennyiben a tartalom nyilvánvalóan művészeti, kreatív, szatirikus, fiktív vagy hasonló mű vagy program részét képezi, az e bekezdésben meghatározott átláthatósági kötelezettségek az ilyen létrehozott vagy manipulált tartalom meglétének megfelelő, a mű megjelenítését vagy élvezetét nem akadályozó közlésére korlátozódnak.*³⁶⁹

Azon túl, hogy a kelleténél tágabban értelmezhető az a kikötés, hogy *„meglétének megfelelő, vagy az élvezetet nem akadályozó”* módon jelöljék a mesterséges tartalom voltát³⁷⁰, e szabályozás valódi hiányossága abból adódik, hogy nem lép fel kellő erővel a szintetikus, mesterségesen generált illetve az ember által létrehozott autentikus tartalmak feltűnésmentes keveredése ellen.

E hiányossághoz kötődik, hogy az 50. cikk (4) bekezdése főszabályként azt írja elő, hogy a *nyilvánosság közérdekű ügyekről való tájékoztatása céljából közzétett szöveget generáló vagy manipuláló MI-rendszer alkalmazóinak közölniük kell, hogy a szöveget mesterségesen hozták létre vagy manipulálták.* Például egy hír- vagy blogoldal, amely nagy nyelvi modellekkel gyártott automatizált cikkeket közöl (folyamatosan frissülő választási eredmények, vagy

³⁶⁷ Kivéve, ha ez a körülményekre és a felhasználási kontextusra figyelemmel, egy észszerűen jól tájékozott, figyelmes és körültekintő természetes személy szempontjából nyilvánvaló tény.

³⁶⁸ MI rendelet 50. cikk (1)

³⁶⁹ MI-rendelet 50. cikk (4)

³⁷⁰ Az (5) bekezdés ezt annyival pontosítja csak, hogy *az említett tájékoztatást legkésőbb az első interakció vagy kitettség alkalmával, egyértelmű és jól megkülönböztethető módon kell az érintett természetes személyek számára nyújtani.*

tőzsdei hírek) főszabályként köteles feltüntetni, hogy szintetikus tartalmat olvas a néző. Azonban a bekezdés második része enyhít ezen a követelményen és azt mondja, *hogy a kötelezettség nem alkalmazandó abban az esetben, ha a közzétett tartalmon emberi felülvizsgálatra vagy szerkesztési ellenőrzésre került sor, és amennyiben a tartalom közzétételéért természetes vagy jogi személy szerkesztői felelősséget visel.*³⁷¹

Itt is megfigyelhető az információ tartalmára összpontosító szabályozási megközelítés. Amennyiben természetes személy felelősséget vállal a közzétételért, a rendelet nem látja szükségességét annak, hogy az olvasó számára biztosítsa a jogot arra, hogy tájékoztassák a tartalom mesterséges eredetéről.

Az előzőekben tárgyalt pszichológiai és technológiai információfeldolgozási torzulás túlságosan sebezhetővé tette a társadalmat a szintetikus tartalmak elterjedésével szemben. A bizalomvesztés és valóságtagadás már napjainkban is megfigyelhető teher. Ha a mindenki számára szabadon elérhető GenMI modellek korábban nem tapasztalt sebeséggel és intenzitással kezdik előállítani a mesterséges tartalmakat (szöveg, kép, és videó formájában egyaránt) ez a folyamat elkerülhetetlenül tovább fog eszkalálódni.

³⁷¹ MI rendelet 50. cikk (4)

6.13 Irodalomjegyzék

- Access Now (2021): Access Now's submission to the European Commission's adoption consultation on the AI Act. 15.o. Elérhető: <https://www.accessnow.org/wp-content/uploads/2021/08/Submission-to-the-European-Commissions-Consultation-on-the-Artificial-Intelligence-Act.pdf> (2022. 05. 09.)
- Alemohammad, S., Casco-Rodriguez, J., Luzi, L., Humayun, A., Babaei, H.R., LeJeune, D., Siahkoohi, A., & Baraniuk, R. (2023): Self-Consuming Generative Models Go MAD. 7-15.o arXiv:2307.01850
- Alignment Research Center (2023): Update on arc's recent eval efforts. Elérhető: <https://metr.org/blog/2023-03-18-update-on-recent-evals> (2023. 03. 18.)
- Allyn, Bobby (2022): Deepfake Video of Zelenskyy Could Be 'Tip of the Iceberg' in Info War, Experts Warn. *NPR*. Elérhető: <https://www.npr.org/2022/03/16/1087062648/deepfake-video-zelenskyy-expertswar-manipulation-ukraine-russia> (2022. 03. 16.)
- Arendt, Hannah (1995): Igazság és politika. In: *Múlt és jövő között. Nyolc gyakorlat a politikai gondolkodás terén.* (ford: Módos Magdolna). Budapest, Osiris–Readers International, 233-251.o.
- Bai, Hui, Jan G. Voelkel, Johannes C. Eichstaedt, and Robb Willer. (2023): Artificial Intelligence Can Persuade Humans on Political Issues. OSF Preprints. <https://doi.org/10.31219/osf.io/stakv>
- Baker, Catherine, Debbie Ging, Maja Brandt Andreassen (2024): Recommending Toxicity: The role of algorithmic recommender functions on YouTube Shorts and TikTok in promoting male supremacist influencers. DCU Centre, Dublin City University. 2-33o. Elérhető: <https://antibullyingcentre.ie/wp-content/uploads/2024/04/DCU-Toxicity-Full-Report.pdf>

- Banasik, J., Crook, J., & Thomas, L. (2003): Sample selection bias in credit scoring models. *Journal of the Operational Research Society*, 54(8), 822–832.o.
<https://doi.org/10.1057/palgrave.jors.2601578>
- Bara Zoltán (2017): Dinamikus árazás az online kereskedelemben – Hogyan lehet hátrányos a fogyasztónak a dinamikus árazás? In: Versenytükör, XIII. évfolyam (2), 4-19.o.
- Barcsi Tamás - Diósi Szabolcs (2022): Hasznos test, dividuum, nyersanyag. A felügyeleti társadalomtól az (ön)ellenőrző és megfigyelési kapitalizmusig. *Magyar filozófiai szemle* 66(3), 190-191.o.
- Barocas, Solon – Selbst, Andrew D. (2016): Big Data’s Disparate Impact. *California Law Review*. 104. 671–732.o
- Bender, E.M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021): On the dangers of stochastic parrots: can language models be too big? *Proceedings of FAccT '21: Conference on Fairness, Accountability, and Transparency*. 610–623.o.
- Bertuzzi, Luca (2023): AI Act: EU countries mull options on fundamental rights, sustainability, workplace use. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/ai-act-eu-countries-mull-options-on-fundamental-rights-sustainability-workplace-use/> (2024. 02. 25.)
- Bertuzzi, Luca (2023): EU’s AI Act negotiations hit the brakes over foundation models. EURACTIV. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/eus-ai-act-negotiations-hit-the-brakes-over-foundation-models/> (2024. 01. 24.)
- Bertuzzi, Luca (2023): France, Germany, Italy push for ‘mandatory self-regulation’ for foundation models in EU’s AI law. EURACTIV. Elérhető: <https://www.euractiv.com/section/artificial-intelligence/news/france-germany-italy-push-for-mandatory-self-regulation-for-foundation-models-in-eus-ai-law> (2024. 01. 20.)
- Blattner, L., és Nelson, S. (2021): How Costly is Noise? Data and Disparities in Consumer Credit. 27-30.o. ArXiv, abs/2105.07554

- Bostrom, Nick (2015): Szuperintelligencia. Utak, veszélyek, stratégiák. (ford. Hidy Mátyás). Budapest, Ad Astra. 89-96.o.
- Bobadilla, J., Ortega, F. Hernando, A. and Gutierrez, A. (2013): Recommender Systems Survey. *Knowledge-Based System*. Vol(46), 112-113.o.
- Bradford, Anu (2020): The Brussels Effect: How the European Union Rules the World. New York, Oxford University Press
- Brendan Nyhan, Jason Reifler (2015): Displacing Misinformation about Events: An Experimental Test of Causal Corrections, *Journal of Experimental Political Science*. Volume 2, Issue 1, 81-93.o.
- Brewer, Jordan, Dhru Patel, Dennie Kim, Alex Murray (2023): Navigating the challenges of generative technologies: Proposing the integration of artificial intelligence and blockchain. *Business Horizons Journal Pre-proof*, 4-7.o.
- Brewster, Jack, Zack Fishman, Elisa Xu (2023): Misinformation Monitor: June 2023. Funding the Next Generation of Content Farms. Elérhető: <https://www.newsguardtech.com/misinformation-monitor/june-2023/>
- Bsharat Sondos Mahmoud, Aidar Myrzakhan, Zhiqiang Shen (2023): Principled Instructions Are All You Need for Questioning LLaMA-1/2, GPT-3.5/4. 5-6.o.
- Burrell, Jenna (2016): How the Machine 'Thinks: Understanding Opacity in Machine Learning Algorithms. *Big Data and Society*. 3/1. 3-6.o.
- Campbell, M., Hoane Jr, A. J., & Hsu, F. H. (2002): Deep Blue. *Artificial Intelligence*, 134(1-2), 59–60.o.
- Castells, Manuel. (2002): Az Internet-galaxis (Gondolatok internetről, üzletről és társadalomról). Budapest, Network Twenty One Hungary kiadó.
- Cheney-Lippold, John (2017): We Are Data: Algorithms and the Making of Our Digital Selves. New York: NYU Press. 88-92.o.

- Chesney, Bobby, Danielle Citron (2019): Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, Vol. 107, 1775-1786.o.
- Cox, Joseph (2024): Google News Is Boosting Garbage AI-Generated Articles. 404media. Elérhető: <https://www.404media.co/google-news-is-boosting-garbage-ai-generated-articles/> (2024. 01. 18.)
- Fantoly Zsanett, Herke Csongor and Szabó Barbara (2023): The role of AI-based systems in negotiated proceedings. *e-Revue Internationale de Droit Pénal*. 2023, A-18, 4.o.
- Collingridge, David (1980): *The Social Control of Technology*. New York. St. Martin's Press
- Cox, Joseph (2023): GPT-4 Hired Unwitting TaskRabbit Worker By Pretending to Be 'Vision-Impaired' Human. VICE. Elérhető: <https://www.vice.com/en/article/jg5ew4/gpt4-hired-unwitting-taskrabbit-worker> (2023. 05. 10.)
- Curreli, Eleonora (2023): ChatGPT Case: How the Italian Data Protection Authority Is Trying To Address AI Risks. MONDAQ. Elérhető: <https://www.mondaq.com/italy/privacy-protection/1317670/chatgpt-case-how-the-italian-data-protection-authority-is-trying-to-address-ai-risks> (2023. 05. 08.)
- Csepeli György (2020): *Ember 2.0 – A mesterséges intelligencia gazdasági és társadalmi hatásai*. Budapest, Kossuth Kiadó
- Danaher J. (2016): The threat of algocracy: Reality, resistance and accommodation. *Philosophy & Technology*. 29 (3), 245-268.o.
- Danaher, J. (2019): The ethics of algorithmic outsourcing in everyday life. In K. Yeung & M. Lodge (szerk.) *Algorithmic Regulation* (91–118o.). Oxford University Press. 107.o. <https://doi.org/10.1093/oso/9780198838494.003.0005>
- Diamond, Larry (2010): Liberation Technology. *Journal of Democracy*. vol. 21, no. 3, 69-83.o.

- Diósi Szabolcs (2022): Behaviorally informed regulations- an emerging trend in modern public policymaking. In Bendes Ákos L. et al. (szerk): III. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar Doktori Iskola. 125-142.o.
- Diósi Szabolcs (2021): The rise of algorithmic decision-making in the age of Big Data. In Bendes Ákos L. et al. (szerk): I-II. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar Doktori Iskola. 150-165.o.
- Diósi Szabolcs (2023): Tiltott gyakorlatok, "elfogadhatatlan kockázat": Az Európai Bizottság rendelettervezete a diszruptív Mesterséges-Intelligencia-technológiák megfékezésére. Európai Jog 23(3), 17-24.o.
- Diósi Szabolcs, Barcsi Tamás (2021): The legacy of disciplinary society – how relevant is Foucault's theory today? Évkönyv - Újvidéki egyetem magyar tannyelvű tanítóképző. XVI. évfolyam, 1. szám, 10-33.o.
- DiResta, R. - Goldstein, J.A. (2024). How Spammers and Scammers Leverage AI-Generated Images on Facebook for Audience Growth. ArXiv, abs/2403.12838.
- Dupré H. Maggie (2023): AI Loses Its Mind After Being Trained on AI-Generated Data. Futurism. Elérhető: <https://futurism.com/ai-trained-ai-generated-data> (2023. 12. 07.)
- Edwards, Benj (2024): OpenAI Can Re-Create Human Voices—but Won't Release the Tech Yet. WIRED. Elérhető: <https://www.wired.com/story/openai-voice-engine-artificial-intelligence-release> (2024. 04. 03.)
- Eli Pariser (2021): The Filter Bubble. What the Internet is Hiding from You New York. The Penguin Press. 112.-113.o.

- Eubanks, Virginia (2018): Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. New York, St Martin's Publishing. 14-39.o.
- Európai Adatvédelmi Testület (2021): Az Európai Adatvédelmi Testület és az európai adatvédelmi biztos 5/2021. sz. közös véleménye. 13-16.o.
- Európai Bizottság (2019): Az emberközpontú mesterséges intelligencia iránti bizalom növelése COM(2019) 168 final
- Európai Bizottság (2020): Az Európai Demokráciára vonatkozó cselekvési tervről. COM(2020) 790 final
- Európai Bizottság: Az Európai Parlament és a Tanács rendelete a digitális szolgáltatások egységes piacáról (digitális szolgáltatásokról szóló jogszabály) és a 2000/ 31/EK irányelv módosításáról. COM(2020)825 final
- Európai Bizottság (2021): A dezinformáció elleni gyakorlati kódex megerősítésére vonatkozó iránymutatás. COM(2021) 262 final,
- Európai Bizottság (2020): Európa digitális jövőjének megtervezése. A Bizottság közleménye az Európai Parlamentnek, a Tanácsnak, az Európai Gazdasági és Szociális Bizottságnak és a Régiók Bizottságának, COM/2020/67, final
- Európai Bizottság (2018): Európai megközelítés az online félretájékoztatás kezelésére. COM(2018) 236 final, 2.1. pont.
- Európai Bizottság (2018): A mesterséges intelligenciáról szóló összehangolt terv COM(2018) 795 final.
- Európai Bizottság (2018): Mesterséges intelligencia Európa számára (A közös európai adattér felé) COM (2018) 237 final. 7-20.o.
- Európai Bizottság (2020): Fehér könyv a mesterséges intelligenciáról – A kiválóság és a bizalom európai megközelítése COM(2020) 65 final. 3-4.o.

- Európai Bizottság, Kommunikációs Főigazgatóság, Leyen, U. (2019): Ambiciózusabb Unió – Programom Európa számára: politikai iránymutatás a hivatalba lépő következő Európai Bizottság számára (2019–2024). Elérhető: <https://data.europa.eu/doi/10.2775/56352>
- Európai Bizottság, Mesterséges intelligenciával foglalkozó magas szintű szakértői csoport (2019): Etikai iránymutatás a megbízható mesterséges intelligenciára vonatkozóan, 19-24.o.
- Európai Gazdasági és Szociális Bizottság (2021): Vélemény - Javaslat európai parlamenti és tanácsi rendeletre a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról INT/940, 4-5.o
- Európai Parlament és a Tanács (EU) 2016/679 rendelete a természetes személyek védelméről a személyes adatok kezelése során és az ilyen adatok szabad áramlásáról, valamint a 95/46/EK irányelv hatályon kívül helyezéséről (Általános Adatvédelmi Rendelet)
- Európai Tanács, Az Európai Tanács ülése (2018. június 28.) – Következtetések, EUCO 9/18 2018, 8.o.
- Európai Tanács (2017): Következtetések, EUCO 14/17, 7.o.
- Európai Tanács 2018: Következtetések. EUCO 9/18, 5. o
- Európai Unió Tanácsa (2022): Javaslat – Az Európai Parlament és a Tanács rendelete a mesterséges intelligenciára vonatkozó harmonizált szabályok (a mesterséges intelligenciáról szóló jogszabály) megállapításáról és egyes uniós jogalkotási aktusok módosításáról – Általános megközelítés. 14954/22
- European Commission (2018): A multi-dimensional approach to disinformation. Report of the independent High level Group on fake news and online disinformation. Luxembourg, Publications Office of the European Union, 22-31.o.
- European Commission (2018): Action Plan against Disinformation. JOIN(2018) 36 final, 5-11.o.

- Fábio Perez, Ian Ribeiro (2022): Ignore previous prompt: Attack techniques for language models. 3-6.o. arXiv.2211.09527
- Fair Trials (2021): Automating Injustice - The Use of Artificial Intelligence & Automated Decision- Making Systems in Criminal Justice. 8-21.o
- Fantoly Zsanett, Herke Csongor, Szabó Barbara (2023): The role of AI-based systems in negotiated proceedings. *e-Revue Internationale de Droit Pénal*. 2023, A-18, 4.o.
- Festinger, L., Riecken, H.W., & Schachter, S. (1956): When Prophecy Fails. University of Minnesota Press
- Festinger, L. (1957): A Theory of Cognitive Dissonance. Stanford University Press.
- Fejes Erzsébet - Futó Iván (2021): Mesterséges intelligencia a közigazgatásban – az érdemi ügyintézés támogatása. *Pénzügyi Szemle*. Különszám 2021/1, 44.o.
- Floridi, Luciano (2020): AI and Its New Winter: from Myths to Realities. *Philosophy & Technology*. 33, 1–3.o. <https://doi.org/10.1007/s13347-020-00396-6>
- FRA Report (2020): European Union Agency for Fundamental Rights (2020): Artificial Intelligence, Big Data and Fundamental Rights Country - Research Estonia 2020.
- Gajanan, Mahita (2017): Kellyanne Conway Defends White House's Falsehoods as 'Alternative Facts'. TIME. Elérhető: <https://time.com/4642689/kellyanne-conway-sean-spicer-donald-trump-alternative-facts/> (2017. 01. 12.)
- Gartner (2012): Gartner Research. The Importance of 'Big Data': A Definition. Elérhető: <https://www.gartner.com/en/documents/2057415>
- Gálík Mihály (2019): A hálózati hírmédia sajátosságai különös tekintettel a visszhangkamra- és a szűrőbuborék-jelenségre. *In Medias Res*, 8, 2. 330–342.o.

- Goldstein A. Josh, Jason Chao, Shelby Grossman, Alex Stamos, Michael Tomz (2024): How persuasive is AI-generated propaganda? *PNAS Nexus*. Volume 3, No 2, 1-7.o.
- Gombos Katalin, Gyuranecz Franciska Zsófia, Krausz Bernadett, Papp Dorottya (2021): A mesterséges intelligencia jogalkalmazási területen való hasznosíthatóságának alapjogi kérdései. In Török Bernát és Zódi Zsolt (szerk): A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 331-332.o.
- Goodfellow, Ian, Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014): Generative Adversarial Nets. 2-8.o. arXiv:1406.2661
- Google (2023): Environmental Report. 49-55.o.
- Gorenz D, Schwarz N (2024): How funny is ChatGPT? A comparison of human- and A.I.-produced jokes. *PLoS ONE* 19(7), 2-10.o. <https://doi.org/10.1371/journal.pone.0305364>
- Guld Ádám (2023): A deepfake és CGI-technológia az influencer marketing szolgálatában: így formálják át a digitális karakterek az ismertségipar működését. In: Aczél, Petra; Veszelszki, Ágnes (szerk.) *Deepfake : a valótlan valóság*. Budapest, Gondolat Kiadó. 189-190.o.
- Grynbaum, M Michael, Ryan Mac (2023): The Times Sues OpenAI and Microsoft Over A.I. Use of Copyrighted Work. *The New York Times*. Elérhető: <https://www.nytimes.com/2023/12/27/business/media/new-york-times-open-ai-microsoft-lawsuit.html> (2023. 12. 27.)
- Gyires-Tóth Bálint (2020): A mélytanulás múltja, jelene és jövője. *Híradástechnika*. 75/1., 23-29.o.
- Hacker, Philipp, Andreas Engel, Marco Mauer (2023): Regulating ChatGPT and Other Large Generative AI Models. In *Proceedings of the Fairness, Accountability, and Transparency (FAccT '23)*. ACM, New York. 1115.o. <https://doi.org/10.1145/3593013.3594067>

- Hagendorf, Thilo (2024): Deception abilities emerged in large language models. PNAS Vol. 121 No. 24, 1-8.o. <https://doi.org/10.1073/pnas.2317967121>
- Haidt, Jonathan, Nick Allen (2020): Scrutinizing the effects of digital technology on mental health. *Nature*. Vol 578, 226-227o.
- Hainsdorf, Clara, Tim Hickman, Sylvia Lorenz, Jenna Rennie (2023): Dawn of the EU's AI Act: political agreement reached on world's first comprehensive horizontal AI regulation. White & Case Tech Newsflash. Elérhető: <https://www.whitecase.com/insight-alert/dawn-eus-ai-act-political-agreement-reached-worlds-first-comprehensive-horizontal-ai> (2024. 01. 30.)
- Hameleers, M, T. E. Powell, T. G. Van Der Meer, L. Bos. (2020): A Picture Paints a Thousand Lies? The Effects and Mechanisms of Multimodal Disinformation and Rebuttals Disseminated via Social Media. *Political Communication*, 37(2):281–301.o.,
- Han, Byung-Chul (2020): Pszichopolitika. (Ford. Csordás Gábor) Budapest, Typotex. 91.o.
- Hassabis, D. (2017): Artificial Intelligence: Chess match of the century. *Nature* 544. 413–414.o. <https://doi.org/10.1038/544413a>
- Harari, Yuval Noah (2018): 21 Lecke a 21. századra. (ford: Torma Péter). Budapest, Animus kiadó. 55-60.o.
- Harari, Y. N. (2023): Yuval Noah Harari argues that AI has hacked the operating system of human civilisation. *The Economist*. <https://www.economist.com/byinvitation/2023/04/28/yuval-noah-harari-argues-that-ai-has-hacked-the-operating-systemof-human-civilisation> (2023. 04. 28.)
- Herke Csongor (2021): A mesterséges intelligencia kriminalisztikai aspektusai. *Belügyi Szemle*. 69/10. 1714–1720.o.
- Héder Mihály (2020): Mesterséges Intelligencia. Filozófiai kérdések, gyakorlati válaszok. Budapest, Gondolat Kör Kiadó

- Hohmann Balázs (2021): A mesterséges intelligencia közigazgatási hatósági eljárásban való alkalmazhatósága a tisztességes eljáráshoz való jog tükrében. In Török Bernát és Zódi Zsolt (szerk.) A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 413.o.
- Houy, C., Hamberg, M. & Fettke, P., (2019): Robotic Process Automation in Public Administrations. In: Räckers, M., Halsbenning, S., Rätz, D., Richter, D. & Schweighofer, E. (Hrsg.), Digitalisierung von Staat und Verwaltung. Bonn: Gesellschaft für Informatik. 62-74.o.
- Howard, Philip N. és Muzammil M. Hussain (2013): Democracy's Fourth Wave? Digital Media and the Arab Spring. Oxford University Press
- Hubert, K. F., Awa, K. N., & Zabelina, D. L. (2024). The current state of artificial intelligence generative language models is more creative than humans on divergent thinking tasks. Scientific reports, 14(1), 3440.
- Hughes, H., & Waismel-Manor, I. (2020): The Macedonian Fake News Industry and the 2016 US Election. PS: Political Science & Politics, 54, 19 – 23.o.
<https://doi.org/10.1017/S1049096520000992>.
- Human Rights Watch (2021): How the EU's Flawed Artificial Intelligence Regulation Endangers the Social Safety Net. 25.o.
- Imperva (2024): Bad Bot Report - 2024. Elérhető:
<https://community.imperva.com/blogs/percy-smith/2024/05/16/impervas-11th-annual-bad-bot-report-2024> (2024. 05. 16.)
- Isaak, J., & Hanna, M. (2018): User Data Privacy: Facebook, Cambridge Analytica, and Privacy Protection. *Computer*. 51, 56-59.o. <https://doi.org/10.1109/MC.2018.3191268>.
- Italie, Hillel (2023): John Grisham, George R.R. Martin and other authors sue OpenAI for copyright infringement. Los Angeles Times. Elérhető: <https://www.latimes.com/world-nation/story/2023-09-20/john-grisham-george-r-r-martin-and-other-authors-sue-openai-for-copyright-infringement> (2023. 09. 20.)

- Jeffrey Dastin (2018): Insight - Amazon scraps secret AI recruiting tool that showed bias against women. REUTERS. Elérhető: <https://www.reuters.com/article/idUSKCN1MK0AG> (2021. 04. 11.)
- Jackie; Raskino, Mark (2008): Mastering the Hype Cycle: How to Choose the Right Innovation at the Right Time. Massachusetts, Harvard Business Publishing
- Josh Taylor (2023): Meta closes nearly 9,000 Facebook and Instagram accounts linked to Chinese ‘Spamouflage’ foreign influence campaign. The Guardian. Elérhető: <https://www.theguardian.com/australia-news/2023/aug/30/meta-facebook-instagram-shuts-down-spamouflage-network-china-foreign-influence> (2023. 08. 30.)
- Kavanagh, Jennifer and Michael D. Rich (2018): Truth Decay: An Initial Exploration of the Diminishing Role of Facts and Analysis in American Public Life. Santa Monica, CA: RAND Corporation. 21-38.o
- Kirchner JH, Ahmad L, Aaronson S, Leike J (2023): OpenAI: New AI classifier for indicating AI-written text. Elérhető: <https://openai.com/blog/new-ai-classifier-for-indicating-ai-written-text> (2023. 05. 20.).
- Kirk, Robert Ishita Mediratta, Christoforos Nalmpantis, Jelena Luketina, Eric Hambro, Edward Grefenstette, Roberta Raileanu (2024): Understanding the Effects of RLHF on LLM Generalisation and Diversity. 4-10.o. arXiv:2310.06452
- Klein, Daniel (2018): Mighty mouse. *MIT Technology Review*. Elérhető: <https://www.technologyreview.com/2018/12/19/138508/mighty-mouse>
- Kollár Csaba (2020): Kína és a társadalmi kredit rendszere. *Hadtudomány: A magyar hadtudományi társaság folyóirata*. 30:2, 79-97.o.
- Koltay András (2021): Előszó. In Török Bernát és Zódi Zsolt (szerk) A mesterséges intelligencia szabályozási kihívásai. Budapest, Ludovika Egyetemi kiadó. 11.o.

- Krekó Péter, Molnár Csaba (2023): Tényrelativizmus és a hírforrásokkal szembeni elbizonytalanodás a magyar közvéleményben. Lakmusz-HDMO. 5.o. Elérhető: https://politicalcapital.hu/pc-admin/source/documents/hdmo_pc_tanulmany_2_tenyrelativizmus_kozvelemeny_20231130.pdf
- Krekó Péter (2021): Tömegparanoia 2.0 - Összeesküvés-elméletek, álhírek és dezinformáció. Budapest, Athenaeum Kiadó. 22.o.
- Kreps, Sarah & Doug Kriner (2023): How AI Threatens Democracy. *Journal of Democracy*. Volume 34(4), 123-125.o.
- Kruger, Justin & Dunning, David. (1999): Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121–1134.o.
- LeCun, Yann, Yoshua Bengio és Geoffrey Hinton (2015): Deep learning. *Nature* 521 (75539), 436-444.o.
- Lemley, A. Mark - Bryan Casey (2020): Fair learning. *Texas Law Review*. 99, 743.o.
- Leviston, Z., I. Walker, and S. Morwinski (2013): Your opinion on climate change might not be as common as you think. *Nature Climate Change*, 3:334–337.o.
- Lim, N., Kuan, M., Pu, M., Lim, M., & Chong, C. (2022). Metamorphic Testing-based Adversarial Attack to Fool Deepfake Detectors. *2022 26th International Conference on Pattern Recognition (ICPR)*, 2503. <https://doi.org/10.1109/ICPR56361.2022.9956543>.
- Liu, Chuncheng (2019): Multiple Social Credit Systems in China. *Economic Sociology: The European Electronic Newsletter*, Vol 21 (1), 22–32.o.
- Liu, N.F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2023): Lost in the Middle: How Language Models Use Long Contexts. *Transactions of the Association for Computational Linguistics*, 12, 157-173.o.
- Luccioni, A. S., Jernite, Y., Strubell, E. (2023): Power hungry processing: Watts driving the cost of ai deployment? 13-14o. arXiv preprint arXiv:2311.16863

- Mezei Kitti, Bán-Forgács Nóra (2022): Discrimination in the Age of Algorithms. In: Human Rights as a Guarantee of Smart, Sustainable and Inclusive Growth. Budapest; Józsefów, Milton Friedman University; Alcide De Gasperi University of Euroregional Economy, 73-80.o.
- Maciejewski, Mariusz (2016): To do more, better, faster and more cheaply: using big data in public administration. *International Review of Administrative Sciences*. 83, 120–135.o.
- Mayer, J., & Mitchell, J. C. (2012): Third-Party Web Tracking: Policy and Technology. *IEEE Symposium on Security and Privacy*
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1955): *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*. Az eredeti szöveg angol nyelven elérhető: <https://www.formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>
- McNulty, Eileen (2014): Understanding Big Data: The Seven V's. DataConomy. Elérhető: <https://dataconomy.com/2014/05/22/seven-vs-big-data> (2014. 05. 22.)
- Meltwater & We Are Social (2024): Digital 2024 Global Overview Report. Elérhető: <https://datareportal.com/reports/digital-2024-global-overview-report>
- Meta, Bickert, Monika (2024): Our Approach to Labeling AI-Generated Content and Manipulated Media. META Newsroom. Elérhető: <https://about.fb.com/news/2024/04/metas-approach-to-labeling-ai-generated-content-and-manipulated-media/> (2024. 04. 05.)
- Michal Kosinski, D. Stillwell, and T. Graepel (2013): Private traits and attributes are predictable from digital records of human behavior. *Proceedings of the National Academy of Sciences*. 110:5802–5805.o.
- Misuraca, Gianluca., - van Noordt, C. (2020): Overview of the use and impact of AI in public services in the EU. EUR 30255 EN, Publications Office of the European Union, Luxembourg 45-46.o.
- Mnih, Volodymyr, Kavukcuoglu, K., Silver, D., Rusu, A., Veness, J., Bellemare, M., Graves, A., Riedmiller, M., Fidjeland, A., Ostrovski, G., Petersen, S., Beattie, C., Sadik, A.,

Antonoglou, I., King, H., Kumaran, D., Wierstra, D., Legg, S., & Hassabis, D., (2015): Human-level control through deep reinforcement learning. *Nature*, 529-533.o.

- Mok, Aaron (2023): A disturbing AI phenomenon could completely upend the internet as we know it. Business Insider. Elérhető: <https://www.businessinsider.com/ai-model-collapse-threatens-to-break-internet-2023-8> (2023. 08. 30.)
- Molnar, Christoph (2020): Interpretable Machine Learning: A Guide for Making Black Box Models Explainable. Leanpub. Elérhető: <https://christophm.github.io/interpretable-ml-book/>
- Moravec, Hans (1988): Mind Children. Cambridge, Harvard University Press. 15-16.o.
- Natasha Lomas (2023): Big Tech's divisive 'personalization' attracts fresh call for profiling-based content feeds to be off by default in EU. TechCrunch. Elérhető: <https://techcrunch.com/2023/12/20/dsa-recommender-systems/>
- N. C. Köbis and L. Mossink (2021): Artificial intelligence versus Maya Angelou: Experimental evidence that people cannot differentiate AI-generated from human-written poetry. *Computers in Human Behavior*, vol.114(2):106553 doi: 10.1016/j.chb.2020.106553
- Németh Tamás (2023): Tilesch György: Sokan még mindig a Transformers-filmek szintjén kezelik a mesterséges intelligenciát. EconomX. Elérhető: <https://www.economx.hu/gazdasag/ai-summit-2023-tilesch-gyorgy-mesterseges-intelligencia-kkv-k-munkaltatok.777160.html> (2023. 09. 11.)
- Nicholas Carlini, Matthew Jagielski, Christopher A Choquette-Choo, Daniel Paleka, Will Pearce, Hyrum Anderson, Andreas Terzis, Kurt Thomas, Florian Tramèr (2023): Poisoning web-scale training datasets is practical.4 -13.o. arXiv:2302.10149.
- Nickerson, Raymond S. (1998): Confirmation bias: A ubiquitous phenomenon in many guises. *Review of General Psychology* 2 (2), 175–220.o.
- Nils J. Nilsson (2009): The Quest for Artificial Intelligence. Cambridge University Press. 13.o.

- O'Hara, K. (2020): Explainable AI and the philosophy and practice of explanation. *Computer Law & Security Review*, Vol.39. 3-6.o.
- OpenAI (2024): Navigating the Challenges and Opportunities of Synthetic Voices. Blog, Product Announcements. Elérhető: <https://openai.com/blog/navigating-the-challenges-and-opportunities-of-synthetic-voices> (2024. 04. 02.)
- OpenAI (2023): GPT-4 Technical Report. 55.o. arXiv:2303.08774v6
- O'Reilly, T. (2005): What Is Web 2.0: Design Patterns and Business Models for the Next Generation of Software. *Communications & Strategies*, No.65 (1), 17-37.o
- Oswald et al (2018): Algorithmic risk assessment policing models: lessons from the Durham HART model and 'Experimental' proportionality. *Information & Communications Technology Law*. 27(2), 223-250.o.
- Pataki Gábor – Szőke Gergely László (2017): Az online személyiségprofilok jelentősége – régi és új kihívások. *Infokommunikáció és jog*. 2017/2., 63–70.o.
- Park A. Lee (2019): Injustice Ex Machina: Predictive Algorithms in Criminal Sentencing. *UCLA Law Review*. Elérhető: <https://www.uclalawreview.org/injustice-ex-machina-predictive-algorithms-in-criminal-sentencing/> (2021. 07. 22.)
- Pasquale, Frank. (2015): *The Black Box Society: The Secret Algorithms That Control Money and Information*. Harvard University Press.
- Peterson, Jake (2024): AI Companies Are Running Out of Internet. LifeHacker. Elérhető: <https://lifehacker.com/tech/ai-is-running-out-of-internet> (2024. 04. 03.)
- Ranerup, Agneta & Henriksen, Z. H (2019): Value positions viewed through the lens of automated decision-making: The case of social services. *Government Information Quarterly*. Volume 36(4) 2-3.o.

- Ranerup, Agneta, Henriksen, Z. H. (2022): Digital Discretion: Unpacking Human and Technological Agency in Automated Decision Making in Sweden's Social Services. *Social Science Computer Review*. 40(2), 445-461.o.
- Régiók Európai Bizottsága (2021): Vélemény. A mesterséges intelligenciával kapcsolatos európai megközelítés – A mesterséges intelligenciáról szóló jogszabály. SEDEC-VII/022, 13.o.
- Robertson, C. E., Pröllochs, N., Schwarzenegger, K., Pärnamets, P., Van Bavel, J. J., & Feuerriegel, S. (2023): Negativity drives online news consumption. *Nature human behaviour*, 7(5), 812–822.o.
- Rossen, Jake (2023): Please Tell Me Your Problem': Remembering ELIZA, the Pioneering '60s Chatbot. *Mental Floss*. Elérhető: <https://www.mentalfloss.com/posts/eliza-chatbot-history> (2023. 02. 14.)
- Rudin, C. (2019): Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nat Mach Intell*. (1), 206–215.o
- Russel, Josh (2023): Sanctions ordered for lawyers who relied on ChatGPT artificial intelligence to prepare court brief. *Courthouse News Service*. Elérhető: <https://www.courthousenews.com/sanctions-ordered-for-lawyers-who-relied-on-chatgpt-artificial-intelligence-to-prepare-court-brief> (2023. 06. 22.)
- Russell, Stuart J. – Norvig, Peter (2010): *Artificial Intelligence: A Modern Approach*. Third Edition. Essex, Pearson. 2-16.o.
- Sadasivan, V., Kumar, A., Balasubramanian, S., Wang, W., & Feizi, S. (2023): Can AI-Generated Text be Reliably Detected?. *ArXiv*, abs/2303.11156;
- Satariano, Adam (2024): France Fines Google Amid A.I. Dispute With News Media. *The New York Times*. Elérhető: https://www.nytimes.com/2024/03/20/business/france-google-fine.htmlutm_source=www.mindstream.news&utm_medium=newsletter&utm_campaign=france-sues-google (2024. 03. 20.)

- Scheurer, J., Balesni, M. and Hobbhahn, M., (2023): Technical Report: Large Language Models can Strategically Deceive their Users when Put Under Pressure. 2-6.o. arXiv:2311.07590
- Schmidhuber, J. (2015): Deep learning in neural networks: An overview. *Neural Networks*, 61, 85-117.o.
- Selbst, A. (2017): Disparate Impact in Big Data Policing. *Georgia law review*. 52, 3373.o. <https://doi.org/10.2139/SSRN.2819182>.
- Shengnan An, Zexiong Ma, Zeqi Lin, Nanning Zheng, Jian-Guang Lou (2024): Make Your LLM Fully Utilize the Context. 3-10.o. arXiv:2404.16811
- Shumailov, I., Shumaylov, Z., Zhao, Y., Gal, Y., Papernot, N., & Anderson, R. (2023): The Curse of Recursion: Training on Generated Data Makes Models Forget. 3 -6.o. ArXiv, abs/2305.17493.
- Smuha, Nathalie A., Ahmed-Rengers, Emma Harkens, Adam, Li, Wenlong MacLaren, James Piselli, Riccardo, Yeung, Karen (2021): How the EU Can Achieve Legally Trustworthy AI: A Response to the European Commission's Proposal for an Artificial Intelligence Act. <http://dx.doi.org/10.2139/ssrn.3899991>
- Stim, Rich: What Is Fair Use? Stanford Libraries. Elérhető: <https://fairuse.stanford.edu/overview/fair-use/what-is-fair-use>
- Subramaniam, R. (2023). Identifying Text Classification Failures in Multilingual AI-Generated Content. *International Journal of Artificial Intelligence & Applications*. Vol.14(5), 57-63.o. <https://doi.org/10.5121/ijaia.2023.14505>
- Sugumaran, D., John, Y., C, J., Joshi, K., Manikandan, G., & Jakka, G. (2023): Cyber Defence Based on Artificial Intelligence and Neural Network Model in Cybersecurity. 2023 Eighth International Conference on Science Technology Engineering and Mathematics (ICONSTEM), 1-8.o. <https://doi.org/10.1109/ICONSTEM56934.2023.10142590>.

- Suleyman, Mustafa, Michael Bhaskar (2023): A következő hullám. Mesterséges intelligencia, technológia, hatalom és a 21. század legnagyobb kihívása. (ford. Farkas Veronika). Budapest, Magnólia kiadó. 74-76.o.
- Sunstein, Cass R. (2001): Republic.com Princeton NJ: Princeton University Press
- Sunstein, Cass (2009): Republic.com 2.0. New Jersey, Princeton University Press. 11.o.
- Sunstein (2017): #Republic: Divided Democracy in the Age of Social Media. Princeton, Princeton University Press.
- Susser, D., Roessler, B., & Nissenbaum, H. (2019): Technology, autonomy, and manipulation. Internet Policy Review, 8(2), 1–21.o.
- Szőke Gergely László (2012): Az Európai adatvédelmi jog megújítása. Tendenciák és lehetőségek az önszabályozás területén. PTE-ÁJK, 57-58.o.
- Szőke Gergely László (2018): Big Data and Algorithms in the Public Sector and Their Impact on the Transparency of Decision-Making. Central and Eastern European eDem and eGov Days. 301-306.o
- Szőke, Gergely László (2021): A közösségi oldalak szabályozási problémái, 105-120., In: Kis Kelemen, Bence; Mohay, Ágoston (szerk.) A technológiai fejlődés jogi kihívásai: Kézikönyv a jogalkotás és jogalkalmazás számára, PTE ÁJK,
- Tegmark, Max (2017): Élet 3.0. Embernek lenni a mesterséges intelligencia korában (Ford.:Weisz Böbe, Garai Attila). HVG Könyvek, Budapest 69.o
- Thaler, R., Sunstein, C. (2008): Nudge: Improving Decisions About Health, Wealth, and Happiness London, Penguin Books
- Tilesch György – Hatamleh, Omar (2021): Mesterség és Intelligencia. Vegyük a kezünkbe a sorsunkat az MI korában. Budapest, Libri Kiadó

- Tim Wu (2016): *The Attention Merchants: The Epic Scramble to Get Inside Our Heads*. New York, Knopf
- Toffler, Alvin (1980): *The Third Wave*. New York, Bantam Books
- Török Bernát (2022): A szólásszabadság a közösségi platformokon és a Digital Services Act. In: Török Bernát és Zódi Zsolt (szerk): *Az internetes platformok kora*. Budapest, Ludovika Egyetemi Kiadó, 197.o.
- Tufekci, Zeynep. (2017): *Twitter and Tear Gas: The Power and Fragility of Networked Protest*. Yale University Press
- Turing, A. M. (1950): Computing Machinery and Intelligence. *Mind*, 59(236), 433 – 460.o.
- U.S. Government Publishing Office (2019): Report of the Select Committee on Intelligence United States Senate on Russian Active Measures Campaigns and Interference in the 2016 U.S. Election. Senate Report, II. Volume, 116–290o.
- van de Donk, W. B. H. J., & Tops, P. W. (1995): Orwell or Athens? Informatization and the Future of Democracy. In W. B. H. J. van Donk, I. T. M. Snellen, & P. W. Tops (szerk.): *Orwell in Athens: a Perspective on Informatization and Democracy*. IOS Press. 13-32.o.
- Von Neumann, J. (1945): First Draft of a Report on the EDVAC. Az eredeti szöveg beszkenelt formában angolul elérhető:
https://archive.computerhistory.org/resources/text/Knuth_Don_X4100/PDF_index/k-8-pdf/k-8-u2593-Draft-EDVAC.pdf
- von Rueden, L., Mayer, S., Beckh, K., Georgiev, B., Giesselbach, S., Heese, R., Kirsch, B., Pfrommer, J., Pick, A., Ramamurthy, R., Walczak, M., Garcke, J., Bauckhage, C. and Schuecker, J., (2019): Informed Machine Learning – A Taxonomy and Survey of Integrating Prior Knowledge into Learning Systems. *IEEE Transactions on Knowledge and Data Engineering*, 35, 614-633.o.

- Vörös Zoltán (2019): Kína a digitális korszakban - Kínai internet és a Társadalmi Kreditrendszer. In: Dombi Judit – Rimai Dávid (szerk.): Digitális forradalom világunkban: tanulmánykötet. Pécs, Institutio Kiadó. 165-182.o.
- Varga Imre (2020): Algoritmusok és a programozás alapjai. A Debreceni Egyetem Informatikai Karának oktatási tananyaga. Elérhető:
<https://irh.inf.unideb.hu/~vargai/APA/index.html>
- Vaswani, A., Shazeer, N.M., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, L., & Polosukhin, I. (2017): Attention is All you Need. *Advances in Neural Information Processing Systems*. 5998-6008.o.
- Veszelszki Ágnes (2017): Az álhírek extra- és intralingvális jellemzői. Századvég 84: 51–82.o.
- Veszelszki Ágnes (2021): deepFAKEnews: Az információmanipuláció új módszerei. In Balász László (szerk.): Digitális Kommunikáció és Tudatosság. Budapest, Hungarovox kiadó, 96.o.
- Viktor Mayer-Schönberger - Kenneth Cukier (2014): Big data. Forradalmi módszer, amely megváltoztatja munkánkat, gondolkodásunkat és egész életünket. Ford. Dankó Zsolt. Budapest, HVG Könyvek
- Vincent, James (2023): OpenAI sued for defamation after ChatGPT fabricates legal accusations against radio host. The Verge. Elérhető:
<https://www.theverge.com/2023/6/9/23755057/openai-chatgpt-false-information-defamation-lawsuit> (2023. 06. 09.)
- Vykopal, I., Pikuliak, M., Srba, I., Móro, R., Macko, D., & Bielíková, M. (2023). Disinformation Capabilities of Large Language Models. ArXiv, abs/2311.08838.
<https://doi.org/10.48550/arXiv.2311.08838>.
- Wells Fargo Bank (2023): Investor Presentation - 2023. 10.o. Elérhető:
<https://www08.wellsfargomedia.com/assets/pdf/about/investor->

relations/presentations/2023/april-investor-presentation.pdf?utm_source=www.mindstream.news&utm_medium=newsletter&utm_campaign=ai-vs-energy-consumption (2023. 08. 12.)

- Wooley, Samuel C. (2020): Bots and Computational Propaganda: Automation for Communication and Control. In *Social Media and Democracy: State of the Field*. (Szerk:) Persily, Nathaniel and Tucker, Joshua A. Cambridge, Cambridge University Press, 89–110.o.
- World Economic Forum (2024): The Global Risk Report 2024 - 19th Edition. 14-37.o.
- W. Youyou, M. Kosinski, and D. Stillwell (2015): Computer-based personality judgments are more accurate than those made by humans. *Proceedings of the National Academy of Sciences*, 112:1036–1040.o.
- Yeung, Karen (2017): Hypernudge: Big data as a mode of regulation by design. *Information, Communication and Society*, Vol.20(1), 118-136.o.
- Yi Liu, Gelei Deng, Zhengzi Xu, Yuekang Li, Yaowen Zheng, Ying Zhang, Lida Zhao, Tianwei Zhang, Kailong Wang, and Yang Liu (2023): Jailbreaking ChatGPT via prompt engineering: An empirical study. 4-9.o. arXiv:2305.13860v2
- Yuval Noah Harari (2018): 21 Lecke a 21. századra. (ford: Torma Péter). Budapest, Animus kiadó. 55-60.o.
- Zhang, Yue, Yafu Li, Leyang Cui, Deng Cai, Lemao Liu, Tingchen Fu, Xinting Huang, Enbo Zhao, Yu Zhang, Yulong Chen, Longyue Wang, Anh Tuan Luu, Wei Bi, Freda Shi, Shuming Shi (2023): Siren's Song in the AI Ocean: A Survey on Hallucination in Large Language Models. 2-6. o. <https://doi.org/10.48550/arXiv.2309.01219>
- Zódi Zsolt (2023): A Collingridge dilemma működés közben. Hogyan szabályozhatunk valamit, amit nem is ismerünk? Ludovika. Elérhető: <https://www.ludovika.hu/blogok/itkiblog/2023/07/31/a-collingridge-dilemma-mukodes-kozben/> (2023. 11. 27.)

- Zódi Zsolt (2017): Jog és jogtudomány a Big Data korában. *Állam- és Jogtudomány 2017/1.*
- Zódi Zsolt (2018): Platformok, robotok és a jog. Új szabályozási kihívások az információs társadalomban. Budapest, Gondolat Kiadó. 234.o.

6.14 Publikációs jegyzék

- Diósi Szabolcs (2023): Adatvezérelt döntések, előrejelző algoritmusok, Mesterséges Intelligencia: Technológiai innováció a köz szolgálatában? In: Barcsi, Tamás (szerk.); Csefkó, Ferenc (szerk.); Diósi, Szabolcs (szerk.): HCS 70. Ünnepi írások Horváth Csaba ny. egyetemi docens születésnapjára. Pécs, Jövő Közigazgatásáért Alapítvány, PTE ÁJK, 74-87.o.
- Diósi Szabolcs (2022): Behaviorally informed regulations- an emerging trend in modern public policymaking. In Bendes Ákos L. et al. (szerk.): III. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar Doktori Iskola, 125-142.o.
- Diósi Szabolcs (2020): Big Data and mental privacy. In: Barna, Boglárka Johanna; Kovács, Petra; Molnár, Dóra; Pató, Viktória Lilla (szerk.): XXIII. Tavaszi Szél Konferencia 2020. Absztraktkötet: MI és a tudomány jövője. Budapest, DOSZ, 522.o.
- Diósi Szabolcs (2022): Facing 'unacceptable risk'. In: Molnár, Dániel; Molnár, Dóra (szerk.) XXV. Tavaszi Szél Konferencia 2022: Absztraktkötet. Budapest, DOSZ, 126.o.
- Diósi Szabolcs (2022): Great opportunities, greater threats: The unfolding history of data-driven social scoring In: Berek, Patrícia; Fodor, Krisztina Dóra (szerk.): Ab ovo usque ad mala – Selected Studies from the "Destinies and Processes" conference. Pécs, Történelmi Ismeretterjesztő Tudás- és Emléktár Alapítvány, 123-136.o.
- Diósi Szabolcs (2023): Hogyan (ne) pontozzuk állampolgárainkat? Ludovika. Kormányzás és Tudomány Blog. Elérhető: <https://www.ludovika.hu/blogok/kormblog/2023/07/14/hogyan-ne-pontozzuk-allampolgarainkat/> (2023. 07.14.)
- Diósi Szabolcs (2019): How algorithm-controlled societies could affect individual autonomy. In: Bódog, Ferenc; Csiszár, Beáta (szerk.): 8th Interdisciplinary Doctoral Conference 2019: Book of Abstract. Pécs, PTE Doktorandusz Önkormányzat. 45.o.
- Diósi Szabolcs (2023): Latest Amendments to the European Union's Artificial Intelligence Act: The Emergence of General-Purpose and Generative AI Technologies. In: Ivanova, Mariana; Dragica, Odzaklieska; Rasim, Yilmaz (szerk.) XX. International Balkan and Near Eastern Congress Series on Economics, Business and Management: Proceedings. Sofia, University St. Kliment Ohridski, Faculty of Economics and Business Administration, 687.o.

- Diósi Szabolcs (2021): Personalized paternalism - present-day analysis of an age-old idea. In: Kajos, Luca Fanni (szerk.); Bali, Cintia (szerk.); Preisz, Zsolt (szerk.); Polgár, Petra [Polgár, Petra Ibolya (szerk.); Glázer-Kniesz, Adrienn (szerk.); Tislér, Ádám [szerk.]; Szabó, Rebeka (szerk.): 10th Jubilee Interdisciplinary Doctoral Conference : Book of Abstracts 241.o.
- Diósi Szabolcs, Barcsi Tamás (2021): The legacy of disciplinary society – how relevant is Foucault’s theory today? Évkönyv - Újvidéki egyetem magyar tannyelvű tanítóképző. XVI. évfolyam, 1. szám, 10-33.o
- Diósi Szabolcs (2021): The rise of algorithmic decision-making in the age of Big Data. In Bendes Ákos L. et al. (szerk.): I-II. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar Doktori Iskola. 150-165.o.
- Diósi Szabolcs (2023): Tiltott gyakorlatok, "elfogadhatatlan kockázat": Az Európai Bizottság rendelettervezete a diszruptív Mesterséges-Intelligencia-technológiák megfékezésére. Európai Jog 23(3), 17-24.o.
- Diósi, Szabolcs (2023): Trustworthy AI in Public Administration - the EU perspective on regulating disruptive technologies. In: Ivanova, Mariana; Nikoloski, Dimitar; Yilmaz, Rasim (szerk.): Proceedings of XVII. International Balkan and Near Eastern Social Sciences Congress Series on Economics. Mitrovica, University of Isa Boletini, 687.o.
- Barcsi Tamás - Diósi Szabolcs (2022): Hasznos test, divídium, nyersanyag. A felügyeleti társadalomtól az (ön)ellenőrző és megfigyelési kapitalizmusig. Magyar filozófiai szemle 66(3), 190-191.o.
- Szabó, Gábor, Diósi, Szabolcs (2022): Ecological debt and sustainable development. In: Simon, Zoltán; Ziegler, Dezső Tamás (szerk.): European Politics - Crises, Fears, and Debates. Paris, L'Harmattan, 89-102.o.

Szerkesztői munkák

- Barcsi Tamás (szerk.); Csefkó, Ferenc (szerk.); Diósi, Szabolcs (szerk.): HCS 70. Ünnepi írások Horváth Csaba ny. egyetemi docens születésnapjára. Pécs, Jövő Közigazgatásáért Alapítvány, PTE ÁJK (2023)

- Barcsi Tamás; Diósi, Szabolcs (szerk.): A válság elméleti vonatkozásai: Tanulmánykötet. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar (2022)
- Barcsi Tamás, Diósi Szabolcs, Mészáros Gábor, Monori Gábor (szerk.): Globális igazságosság, emberi jogok, jogász etika: Szabó Gábor- emlékkötet. Pécs, Pécsi Tudományegyetem Állam- és Jogtudományi Kar (2023)
- Tóth Dávid; Bendes, Ákos László; Diósi, Szabolcs; Gáspár, Zsolt; Projics, Nárcisz; Serbakov, Márton Tibor; Szívós, Alexander Roland (szerk.) I-II. Konferenciakötet : A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, PTE Állam- és Jogtudományi Kar, Doktori Iskola (2021)
- Bendes Ákos László (szerk.) ; Diósi, Szabolcs (szerk.) ; Gáspár, Zsolt (szerk.) ; Gáti, Balázs (szerk.) ; Projics, Nárcisz (szerk.) ; Tóth, Dávid (szerk.): III. Konferenciakötet: A pécsi jogász doktoranduszoknak szervezett konferencia előadásai. Pécs, PTE Állam- és Jogtudományi Kar, Doktori Iskola (2022)