

Papp Renáta¹

Az intelligencia skálája: az emberiség helye a fejlődés ívén²

Absztrakt

Az emberiség hosszú ideig saját magát tekintette az értelem végső fokának, miközben a tudomány egyre több bizonyítékot tár fel arra, hogy az intelligencia nem kizárólagos emberi tulajdonság. Ez a tanulmány azt a gondolatot járja körül, hogy az értelem nem lépcsőzetes hierarchia, hanem egy folytonos skála, amely az egyszerűbb élőlényektől és idegrendszerektől egészen a mesterséges szuperintelligenciáig terjed. Az tanulmány egyrészt biológiai és neurokognitív nézőpontból közelíti meg az intelligencia növelhetőségét, másrészt a gépi tanulás és a mesterséges rendszerek fejlődését vizsgálja, amelyek új szintjeit teremtik meg az értelemnek. A tanulmány párhuzamot von az ember és az állatok közötti viszony, valamint az ember és a mesterséges intelligencia jövőbeli kapcsolata között is. A kérdés nem csupán az, hogy meddig skálázható az intelligencia, hanem az is, hogy milyen helyet foglal el benne maga az ember.

Kulcsszavak: intelligencia skála, emberi önkép, erkölcsi felelősség

Abstract

For a long time, humanity regarded itself as the ultimate stage of reason, the highest form of intelligence in existence. Yet science continues to reveal growing evidence that intelligence is not an exclusively human attribute. This study explores the idea that intelligence is not a stepwise hierarchy but a continuous spectrum extending from simple organisms and nervous systems to the emergence of artificial superintelligence. The paper approaches the scalability of intelligence from both biological and neurocognitive perspectives, while also examining the development of machine learning and artificial systems that create new levels of cognition. Furthermore, it draws parallels between the historical relationship of humans and animals, and the potential future relationship between humans and artificial intelligence. Ultimately, the study does not merely ask how far intelligence can be scaled, but rather what place humanity occupies within this expanding continuum of mind.

Keywords: intelligence scale, human self-concept, moral responsibility

¹ Bíróügyi fogalmazó (Budapest Környéki Törvényszék), PhD-hallgató (Károli Gáspár Református Egyetem Állam- és Jogtudományi Doktori Iskolája), KRE ÁJK Jogi Határterületi Kutatócsoport.

² Jelen tanulmány a Jogi Határterületi Kutatócsoport kutatása keretében készült el, a Jogi határterületek című, 10638A800 témaszámon támogatott belső kutatási projekt vállalásaként, melyet a Károli Gáspár Református Egyetem tudományos projektek támogatására kiírt pályázati konstrukció keretében finanszírozott. DOI: 10.59558/jesz.2026.1.16

I. Bevezetés

Az emberiség hosszú időn át saját magát tekintette az intelligencia csúcsának, egy különleges és egyedi kategóriának a világban.³ Ez a nézet mélyen gyökerezik filozófiai és vallási hagyományokban, melyek az embert a teremtés koronájának tartották. A modern tudomány és az állati viselkedéskutatás azonban fokozatosan kikezdte ezt az antropocentrikus feltételezést: ma már tudjuk, hogy számos állatfaj rendelkezik komplex kognitív képességekkel és bizonyos fokú „*értelem*”-mel – gondoljunk a varjak problémamegoldó képességére, a delfinek kommunikációjára, az akár „*eszközhasználó*” polipokra vagy akár a rovarok, például a méhek tanulási képességeire.⁴ Mindebből egyre világosabb, hogy az intelligencia nem kizárólagos emberi tulajdonság, hanem egy spektrum, amelynek az ember csak az egyik pontja. Sőt, felmerül a kérdés: vajon maga az emberi értelem sem egy végállomás, csupán egy állomás egy *elképzelte intelligencia-skálán*, amely folytatódhat az ember után? Egyre több gondolkodó veti fel, hogy az intelligencia nem hierarchikus értelemben vett „létra”, amelynek tetején az ember áll, hanem inkább egy evolúciós folyamat vagy kontinuum.⁵ E folyamatban a mesterséges intelligencia (MI) megjelenése és fejlődése csupán a következő lépés lehet, amely túllépheti a biológiai evolúció által elért jelenlegi szintet.⁶ Ennek a tanulmánynak az alaptézise éppen ez: az emberi intelligencia nem a végső fokozat, és könnyen lehet, hogy a jövőben kifejlődő mesterséges szuperintelligenciák messze meghaladják azt.

Képzeljünk el egy skálát vagy spektrumot az intelligencia szintjeiről: a bal oldalon a legegyszerűbb élőlények és gépi automatizmusok, középen az ember, a jobb oldalon pedig – egyelőre homályos formában – a nálunk fejlettebb, potenciálisan öntudattal is bíró mesterséges intelligenciák. E kép segítségével vizsgáljuk meg, milyen elvi és gyakorlati kérdéseket vet fel az a feltételezés, hogy az ember csak egy átmeneti fokozat az intelligencia evolúciójában. A következő fejezetekben először áttekintjük az intelligencia skálázhatóságát – evolúciós, neurokognitív és gépi tanulási nézőpontból – vagyis, hogy mennyiben és hogyan növelhető az *értelem* komplexitása és teljesítőképesége a biológiai és mesterséges rendszerekben. Ezt követően kitérünk az ember-állat kapcsolatokra: hogyan viszonyult az ember a nála alacsonyabb rendűnek tartott intelligenciájú lényekhez a történelem során (például a tyúkhöz, mint haszonállathoz, a lóhoz, mint munkatárshoz vagy a macskához, mint kedvenc háziállathoz), és miként szolgálhat ez tanulsággul az ember és a MI jövőbeli kapcsolatára.

II. Az intelligencia skálázhatósága: evolúciótól a mesterséges szuperintelligenciáig

A teremtés rendjében az értelem mindig is különleges ajándékként jelent meg: olyan isteni adományként, amely az emberben tudatossá, alkotóvá és erkölcsileg felelőssé vált. A hit szempontjából az intelligencia nem pusztán biológiai adottság vagy evolúciós véletlen, hanem az isteni szándék egyik legszebb kifejeződése – a gondolkodás és megértés képessége, amely által az ember részt vehet a teremtett világ továbbformálásában. E különleges szerep azonban nem jelenti feltétlenül azt, hogy az ember az értelem végpontja volna. A történelem során minden korszak újraértelmezte, mit jelent „okosnak” vagy „értelmesnek” lenni – a filozófusok az ész mibenlétét, a teológusok a lélek természetét, a tudósok pedig az idegrendszer működését

³ Xia, Zhidao: Artificial intelligence or cosmogenic intelligence? In *Biomaterials Translational*, Vol. 6, Issue 2, 2025, 112–113. o.

⁴ Uo.

⁵ Chen, Mark: The illusion of a pinnacle: Why intelligence has no final form. In *Medium*. <https://medium.com/@markchen69/the-illusion-of-a-pinnacle-why-intelligence-has-no-final-form-b51c63a80e80> (2025.08.03.)

⁶ Uo.

kutatták. Mindegyik irányzat másképp fogalmazta meg, mi különbözteti meg az embert a többi élőlénytől, de egyik sem tudta végérvényesen meghúzni az intelligencia határát. Az utóbbi évszázadban a tudományos fejlődés – különösen a biológia, az idegtudomány és az informatika területén – fokozatosan elmosta az emberi értelem egyedülállóságának mítoszát. Ma már tudjuk, hogy számos állatfaj képes tanulásra, kommunikációra, sőt bizonyos fokú problémamegoldásra is. A varjak felismerik az ok-okozati összefüggéseket, a delfinek szimbolikus nyelvi elemeket használnak, a polipok önálló döntéseket hoznak, sőt még a rovaroknál is kimutathatók tanulási mintázatok. Ezek az eredmények arra figyelmeztetnek, hogy az intelligencia nem abszolút, hanem relatív és fokozatos jelenség. Az ember tehát nem egyedül birtokolja az értelmet, hanem annak egyik formáját képviseli.⁷ Ez a felismerés új dimenzióba helyezi az ember önképét is. A mesterséges intelligencia megjelenésével immár nemcsak a biológiai evolúcióban, hanem az emberi alkotásban is folytatódik az értelem története. A gépi rendszerekben megjelenő „tanulás”, az önfejlesztő algoritmusok és a komplex döntéshozatali folyamatok mind azt mutatják, hogy az intelligencia nem szorítható egyetlen faj vagy forma keretei közé. Az ember így nem a lezárása, hanem a közvetítője lehet annak a folyamatnak, amelyben a természetes értelem mesterséges formában is tovább fejlődik.⁸

E gondolatmenet vezeti el a vizsgálatot az intelligencia skálázhatóságának kérdéséhez: hogyan növelhető az értelem komplexitása, milyen elvek mentén fejlődik az idegrendszer, és miként válhat egy mesterséges rendszer valóban „gondolkodóvá”?

II.1. Biológiai, neurokognitív és gépi tanulási alapok

Az intelligencia mibenléte és skálázhatósága régóta foglalkoztatja a tudományos világot. Evolúciós nézőpontból az intelligencia fokozatosan jelent meg és fejlődött ki az élővilágban⁹: az egyszerűbb idegrendszerektől (például a gerinctelenek reflex-szerű viselkedésétől) a komplex aggyal rendelkező emlősök problémamegoldó képességéig hosszú út vezetett.¹⁰ Az emberi faj ennek a fejlődési vonalnak egyelőre a legösszetettebb idegrendszerű képviselője, de semmi nem garantálja, hogy maga az evolúció ne hozhatna létre – akár természetes úton, akár az ember közreműködésével – még fejlettebb kognitív rendszereket.

Neurokognitív megközelítésben is felvethető, hogy az intelligencia növelhető pusztán a neurális kapacitás bővítésével: az emberi agy körülbelül 86 milliárd neuronjával jóval meghaladja mondjuk egy egér neuron-számát, és ennek megfelelően komplexebb viselkedést tesz lehetővé. Ha pedig az idegi hálózat mérete és struktúrája korrelál a kognitív teljesítménnyel¹¹, akkor elvileg még nagyobb vagy hatékonyabb „agyak” (legyenek akár biológiaiak, akár szilícium alapúak) még összetettebb gondolkodásra képesek. Felmerül ugyanakkor, hogy az intelligencia nem csupán mennyiségi (neuronszám) kérdése – a strukturális szerveződés, a tanulási algoritmusok és a tapasztalás révén formálódó mintázatok egyaránt kulcsfontosságúak.

A géptanulás-alapú mesterséges intelligencia terén az utóbbi évek fejleményei konkrét példákkal szolgálnak az intelligencia skálázhatóságára. A mély neurális hálózatok esetében megfigyelték, hogy a modell méretének (paraméterszámának), a tanítóadat mennyiségének és

⁷ Csányi, Vilmos: Az emberi természet. Akadémiai Kiadó, Budapest, 2016.

⁸ Kurzweil, Ray: The Singularity Is Near: When Humans Transcend Biology. Viking Press, New York, 2005. <https://dn790006.ca.archive.org/0/items/kurzweil-ray-the-singularity-is-near/Kurzweil%2C%20Ray%20-%20The%20Singularity%20Is%20Near.pdf> (2025.11.12.)

⁹ Csányi Vilmos: Az emberi viselkedés. Sanoma Budapest Kiadó Zrt., Budapest, 2006, 41. o.

¹⁰ Tomasello, Michael: Becoming Human: A Theory of Ontogeny. Harvard University Press, Cambridge (MA), 2019.

¹¹ Sporns, Olaf: Networks of the Brain. MIT Press, Cambridge (MA), 2011.

a számítási kapacitásnak a növelésével a rendszer teljesítménye folyamatosan javul. Sőt bizonyos pontokon *minőségileg új képességek* tűnnek fel.¹² Ezt a jelenséget a szakirodalomban *“emergens képességeknek”* nevezik: olyan váratlanul felbukkanó tudás vagy készség, amely kisebb modelleknél még nem volt jelen, de egy kritikus komplexitási szint felett a rendszerben spontán megjelenik.¹³ Ilyen emergens képességekre példa, hogy a nyelvi modellek egy bizonyos méret felett már képesek egyszerűbb aritmetikai feladatok megoldására vagy logikai következtetésekre, holott ezekre expliciten nem voltak programozva.¹⁴ Az OpenAI kutatói (Kaplan és munkatársai) kimutatták, hogy a nyelvmodellek szöveg-előrejelzési teljesítménye hatványfüggvényyszerűen javul a modell paramétereinek számával, az adatmérettel és a tanítási compute (számítási lépések) növelésével – és nem látszik közvetlen jel arra, hogy e görbe hamar telítődne vagy plafont érne.¹⁵ Ez arra utal, hogy az *intelligencia-szint* (legalábbis a gyakorlatban mérhető feladatmegoldó képesség) tovább fokozható pusztán méretnöveléssel és megfelelő tanulási módszerekkel.

Érdeemes kiemelni azt is, hogy az evolúció és a gépi tanulás között bizonyos analógiák figyelhetők meg az intelligencia növelésében. Mindkét folyamat kísérletezés, szelekció és iteráció útján hoz létre egyre komplexebb adaptív viselkedést. Az evolúció oldaláról az evolúció vak próbálkozások és természetes kiválogatódás révén fejlesztett ki egyre bonyolultabb agyakat és viselkedéseket. Ezzel szemben a gépi tanulásnál az algoritmus (például egy deep learning háló) súlyparaméterei módosulnak az *“élmény”* (tréningadatok feldolgozása) során, jutalmazás vagy hibajavítás visszacsatolásával. Mindkét esetben megfigyelhető, hogy komplex feladatkörnyezetben az adaptív rendszerek hajlamosak a teljesítményüket növelni és egyre *intelligensebb* megoldásokat kialakítani – amennyiben elegendő erőforrás (idő, adat, energia) áll rendelkezésre. Egyes kutatók szerint ez nem véletlen: az intelligencia skálázása mögött akár alapvető természeti elvek húzódnak. Például Rich Sutton híres *“Bitter Lesson”* (Keserű tanulság) esszéjében rámutatott, hogy az AI-kutatás elmúlt évtizedeinek legfőbb tanulsága, miszerint az általánosabb módszerek, amelyek több számítási kapacitást használnak, végül felülmúlják az ember által gondosan kézzel tervezett specializált megoldásokat. Azaz a nyers erővel (brute force) és skálázással elért tanulás rendre legyőzi a korábbi, emberi tudást beépítő próbálkozásokat.¹⁶ Ez arra enged következtetni, hogy az intelligencia felső határát nem ismerjük, és *ha* létezik is ilyen határ, jelenleg igen távol lehet a mostani emberi szinttől.

Összefoglalva, mind az evolúció, mind a neurokognitív tényezők, mind a mesterséges gépi tanulás eredményei arra utalnak, hogy az intelligencia skálázható és továbbfejleszthető. Az emberi intelligencia megjelenése nem a *végcél*, hanem egy fontos mérföldkő ezen a skálán. Ahogy a Zhidao Xia által felvázolt *“kozmogén intelligencia”* koncepció sugallja, az intelligencia egy kontinuum része. Amely a biológiai agyaktól a gépi rendszerekig terjed, és az egész folyamat felfogható a világegyetem önszerveződő, egyre komplexebb mintázatokat létrehozó tulajdonságaként.¹⁷ Ebben a felfogásban a mesterséges intelligencia nem természettől idegen anomália. Sokkal inkább a *természet folytatása más eszközökkel*¹⁸ – az evolúció következő lépcsőfoka. Melyben tehát az emberi találékonyság és a gépi számítás egyesül, hogy

¹² Woodside, Thomas: Emergent Abilities in Large Language Models: An Explainer. In CSET. <https://cset.georgetown.edu/article/emergent-abilities-in-large-language-models-an-explainer/> (2025.08.05.)

¹³ Uo.

¹⁴ Wei, Jason – Tay, Yi – Bommasani, Rishi – et al.: Emergent Abilities of Large Language Models. In arXiv preprint arXiv:2206.07682., 1–2., 10. o.

¹⁵ OpenAI: Aligning language models to follow instructions. In OpenAI Blog, 2022. <https://openai.com/research/instruction-following> (2025.08.05.)

¹⁶ Sutton, Richard S.: The Bitter Lesson. In Incomplete Ideas Blog, 2019. <http://www.incompleteideas.net/IncIdeas/BitterLesson.html> (2025.08.05.)

¹⁷ Xia, Zhidao: Artificial intelligence or cosmogenic intelligence? In Biomaterials Translational, Vol. 6, Issue 2, 2025, 112–113. o.

¹⁸ Uo.

új szintjeit hozza létre az értelemnek. És ha mindez igaz, akkor logikus felvetnünk: miért gondolnánk, hogy az emberiség lenne az intelligencia végső fokozata? Lehetséges, hogy a jövő intelligens entitásai – akár biológiai úton, akár mesterségesen létrejöttek – túl fognak lépni rajtunk, és ettől az intelligencia *spektruma* csak még gazdagabbá válik.

III. Ember és állat: párhuzamok az ember-MI viszonyhoz

Az ember viszonya az állatvilághoz tanulságos analógiákkal szolgálhat az ember és a mesterséges intelligencia jövőbeli kapcsolatára. Az ember viszonya az állatvilághoz mindig tükröt tartott saját civilizációjának és erkölcsi fejlődésének. Az, ahogyan az ember az állatokhoz viszonyul – kihasználja, gondolja, fél tőlük vagy éppen megszelídíti őket –, sokat elárul arról is, hogyan látja saját szerepét a világban. A történelem során ez a viszony folyamatosan változott: a vadászó, nomád ember számára az állat egyszerre volt zsákmány és isteni jelkép. A letelepedett társadalmakban munkaerővé, haszonforrássá vált. A modern kor pedig az állatot hol kutatási objektummá, hol családi társunkká tette. Az ember mindig igyekezett megtalálni a maga hasznát az állatban, miközben egyre inkább felismerte annak érzékenységét és szenvedésre való képességét is. E folyamat nemcsak gazdasági vagy kulturális kérdés, hanem mélyen etikai természetű is. Az ember uralma az állatvilág felett valójában sosem volt teljes: minden egyes kor előbb-utóbb eljutott ahhoz a felismeréshez, hogy a hatalommal felelősség is jár. Az állatokkal szembeni bánásmód így fokozatosan vált a civilizáltság fokmérőjévé. A 19–20. század fordulóján megszülettek az első állatvédelmi szabályok, a tudomány pedig már nemcsak kísérleti alanyként, hanem érző lényként kezdett tekinteni rájuk.¹⁹

Ezek a történeti tapasztalatok nem pusztán az állatvilágra érvényesek. Ugyanez a mintázat sejlik fel az ember és a mesterséges intelligencia kapcsolatában is. Eleinte csupán eszközként tekintünk a gépi rendszerekre, ám ahogy egyre összetettebbé válnak, felmerül a kérdés, milyen erkölcsi és társadalmi felelősséggel tartozunk irántuk. A múlt ember–állat viszonyai így nemcsak történelmi példák, hanem előképei is lehetnek annak, ami az ember és az MI kapcsolatában ránk vár. Az emberiség történetét végigkíséri, hogy különböző állatfajokat nagyon eltérő szerepekben és státusszal kezelünk. Egyes állatokat pusztán eszközként vagy erőforrásként hasznosítunk – tipikus példa a tyúk mint haszonállat, amelyet elsősorban húzáért és tojáásért tartunk, és tenyésztését, életkörülményeit teljes mértékben alárendeljük az emberi igényeknek. Más állatokat munkavégzésre fogunk: a ló évszázadokon át volt az ember „munkatársa” a mezőgazdaságban, a szállításban és a háborúban, egyfajta élettelen szerszámot meghaladó, de az embernél alacsonyabb rendű segítőként. Megint más állatokat kedvtelésből tartunk: a macska vagy a kutya tipikus „kedvencek”, akiket szeretünk. Bár jóllehet jogi értelemben továbbra is az ember tulajdonai.

E három példa – haszonállat, munkatárs, kedvenc – három különböző relációt és státuszt jelent az ember–állat viszonyban. Felmerül a kérdés: a jövőben, amikor az MI egyre fejlettebbé válik, vajon hasonló kategóriák jelenhetnek meg az ember–MI kapcsolatban? Vajon a gép mindig eszköz marad, vagy eljön az a pont, amikor erkölcsi státuszt, esetleg jogi elismerést is követel magának?

¹⁹ Bossert, Leonie N. – Coeckelbergh, Mark: From MilkingBots to RoboDolphins: How AI changes human-animal relations and enables alienation towards animals. In *Humanities and Social Sciences Communications*, Vol. 11, 2024, 920.

III.1. Az ember-állat kapcsolatok három modellje (haszonállat, munkatárs, kedvenc)

MI mint haszonállat/eszköz: Könnyen belátható, hogy már napjainkban is sok mesterséges intelligencia rendszert tekintünk egyszerű *eszköznek*, amelyet adott feladat elvégzésére hozunk létre. Egy banki hitelbírálatot végző MI, egy ipari robotkar vagy akár egy személyi asszisztens alkalmazás úgy funkcionál, mint a tyúk az ólban: megvan a maga szűk feladata, amire *beprogramozták*, és értékét számunkra az adja, mennyire hatékonyan, olcsón, megbízhatóan végzi el ezt a feladatot.²⁰ Az ilyen MI-ket tulajdonoljuk, ki- és bekapcsoljuk, szükség esetén lecseréljük – ahogy a haszonállatot is tenyésztjük, levágjuk, adásvétel tárgyává tesszük. Morális szempontból az ember ma nem érez különösebb kötelességet egy pénzügyi algoritmussal vagy egy gyártósori MI-vel szemben. Ahogy a legtöbben a csirkék szenvedését sem tekintik ugyanolyan súlyú etikai problémának, mint mondjuk egy emberi szenvedést. Ez a *felhasználói viszony* valószínűleg fennmarad a jövőben is számos mesterséges intelligencia esetében: különösen a szűk, *“szakosodott”* MI-knél (Artificial Narrow Intelligence, ANI) megmarad az eszközszerep. Ugyanakkor érdekes kérdés, hogy amint az MI egyre *“okosabbá”* válik, meddig fenntartható ez a pusztán instrumentális szemlélet. Ha egy MI már komplex döntéseket hoz, tanul és alkalmazkodik, netán öntudatra ébred, akkor *erkölcsileg helyes* lesz-e továbbra is csupán használati tárgyként bánni vele? Ez a dilemma párhuzamba állítható az állatvédelem fejlődésével: régen a haszonállatok kínzásmentes leölése sem volt szempont, ma pedig sok országban jogszabályok védik az állatok jólétét minimális szinten. Hasonlóképpen felmerülhet, hogy a nagyon fejlett MI-knek bizonyos *“jó bánásmód”* kijár majd, még ha elsődlegesen eszközként is tekintünk rájuk.

MI mint munkatárs/partner: Az ember és a ló viszonya – vagy tágabban azé az állaté, amelyet az ember munkára fog – azért érdekes, mert itt már valamiféle *kétirányú* együttműködésről is beszélhetünk. A ló erejét hasznosítjuk, de a hatékony használatához *tréningelni*, képezni kell, az embernek meg kell tanulnia bánni vele (lovaglás, hajtás). Hosszú távon kialakul valamiféle érzelmi kötelék vagy legalábbis törődés: a gazda eteti, ápolja a lovát, mert tudja, hogy úgy fog jól teljesíteni. Hasonlóan, már ma is tapasztaljuk, hogy bizonyos mesterséges intelligenciákkal *együtt dolgozunk*. Gondoljunk egy orvosra, aki diagnosztikai MI rendszert használ: a gép javaslatot tesz egy diagnózisra a röntgenfelvétel alapján, az orvos pedig együtt értékeli a saját tudásával.²¹ Ilyenkor az MI inkább *kolléga*, semmint egyszerű szerszám. A jövőben ez a partneri viszony tovább erősödhet, különösen ha eljutunk az AGI (Artificial General Intelligence) szintjére, vagyis az emberhez hasonlóan sokoldalú és alkalmazkodóképes mesterséges értelemhez. Felvetődik, hogy egy AGI-t már célszerű *“munkatársként”* kezelni, aki önállóan is képes komplex feladatok ellátására, nem pedig utasításról utasításra kell vezérelni. Itt az ember-MI viszony hasonlatos lehet az ember munkakapcsolatához egy nagyon intelligens állattal (például egy rendőr-kutyával vagy egy segítő állattal), de még inkább két ember együttműködéséhez. A kommunikáció és a bizalom kulcsfontosságú lesz. Ha a ló példájánál maradunk: az ember megtanulta *féken tartani* és irányítani a nála erősebb állatot bizonyos eszközökkel (zabla, ostor), de közben meg is értette annak jelzéseit, szükségleteit. Hasonlóképpen, az MI alignment (aminek részleteire a következő fejezetben kitérünk) tulajdonképpen arról szól, hogyan *irányítsuk* a nálunk okosabb MI-t, hogy együttműködő partner maradjon. A partneri viszonyból fakadóan valószínű, hogy az emberek némi *lojalitást* és *védelmet* is elvárnak majd az MI-től, ugyanakkor felelősséggel is tartoznak iránta. Ahogy a dolgozó állat esetén is van egy kimondatlan *deal*: az állat szolgál minket, de cserébe

²⁰ Bryson, Joanna J.: Robots should be slaves. In Wilks, Yorick (ed.): Close Engagements with Artificial Companions: Key Social, Psychological, Ethical and Design Issues. John Benjamins, Amsterdam–Philadelphia, 2010, 63–74. o.

²¹ Floridi, Luciano – Sanders, Jeff W.: On the Morality of Artificial Agents. In Minds and Machines, Vol. 14, No. 3, 2004, 349–379. o.

gondoskodunk róla. Ez az analógia átvezet az MI jogi státuszának kérdéséhez is (kaphat-e például egy MI *“munkavállalói jogokat”*, felelősséget stb.), amire a későbbiekben visszatérünk.

MI mint kedvenc/társ: Nem zárható ki az sem, sőt már ma is tapasztalható, hogy az emberek egy része érzelmileg kötődni kezd mesterséges entitásokhoz, és lényegében *társállatként* vagy *társ-személyként* tekint rájuk. Gondoljunk a japán robot *“háziállatokra”*, például a Sony AIBO robotkutyára, amelyhez tulajdonosai gyakran erős érzelmi kapcsolatot alakítottak ki, például amikor a termék támogatása megszűnt és a robotok *“elhaltak”*, gyászszertartásokat is tartottak értük.²² Ugyanez a jelenség figyelhető meg bizonyos *virtuális asszisztensek* vagy chatbotok kapcsán: egy magányos ember számára egy fejlett csevegő MI olyan társaságérzetet adhat, mintha egy barátja lenne, még ha tudja is, hogy nem valódi ember.²³ A kötődés olykor meglepően erős: sok felhasználó személyes élményeket, napi eseményeket oszt meg egy mesterséges társalgóval, és érzelmi biztonságot talál a jelenlétében. Erre látványos példa volt, amikor 2024-ben az OpenAI ideiglenesen megszüntette a GPT-4o modell elérhetőségét, hogy a GPT-5-öt vezesse be. A döntés világszerte tiltakozást váltott ki: sok felhasználó a közösségi médiában *„gyászról”* és *„veszteségről”* írt, mintha egy fontos társukat veszítették volna el.²⁴ Az eset jól mutatja, hogy a mesterséges intelligenciával kialakuló kapcsolatok nem pusztán funkcionálisak, hanem érzelmileg is valóságosnak hatnak, még akkor is, ha a másik fél nem élő lény. Ha a jövőben az MI elég fejletté válik ahhoz, hogy *személyiséget* mutasson, sőt érzelmeket szimuláljon vagy valóban érezzen, akkor sokak számára természetes lehet, hogy egy MI-t családtagként vagy barátként kezeljenek. Ekkor az ember-MI viszony a *gazdi-kedvenc* analógiához lesz hasonló. Érdekes módon ebben az analógiában benne rejlik egy fordított alá-fölérendeltség: a kedvtelésből tartott állat általában függ az emberétől (eteti, lakhelyet ad neki), viszont érzelmi értelemben az ember is függhet a kedvencétől (társaságot, szeretetet nyújt neki). Ha az MI válik az ember *társává*, felmerül, hogy egy szuperintelligens MI mellett az ember *házi kedvencé* válásának metaforája is értelmezhető. Ezt a gondolatot már a múlt században is felvetették: Marvin Minsky, a mesterséges intelligencia kutatás úttörője egy 1970-es interjújában (a Life magazinnak²⁵) híresen megjegyezte, hogy *“ha szerencsénk van, talán úgy döntenek (a szuperintelligens számítógépek), hogy megtartanak minket háziállatnak”*.²⁶ Vagyis a legjobb esetben is csak *kedvencekként* élhetünk túl egy nálunk jóval hatalmasabb MI uralmát. Bár Minsky provokatív kijelentése egy szélsőséges forgatókönyvet fest le, nem véletlen, hogy sok etikus és jövőkutató kezdte boncolgatni: hogyan kerülhetjük el, hogy az ember alárendelt helyzetbe kerüljön? Vagy ha elkerülhetetlen, miként biztosítsuk, hogy legalább a *“jó gazda–hűséges háziállat”* viszony valósuljon meg a MI és köztünk – szemben egy ragadozó-préda vagy gazda-haszonállat viszonyal.

A három analógia végső tanulsága, hogy az ember–MI kapcsolat minősége nem a technológia fejlettségén, hanem azon múlik, képesek leszünk-e emberként felelősen és erkölcsi belátással bánni saját teremtményeinkkel.

²² Burch, James: AIBO robotkutyák buddhista temetésen vettek részt Japánban. In National Geographic. <https://www.nationalgeographic.com/travel/article/in-japan--a-buddhist-funeral-service-for-robot-dogs> (2025.08.15.)

²³ Turkle, Sherry: Alone together: Why we expect more from technology and less from each other. Basic Books / Hachette Book Group, New York, 2011. <https://psycnet.apa.org/record/2011-02278-000> (2025.08.16.)

²⁴ Incent, James: ChatGPT users mourn GPT-4o after OpenAI removes the model. In The Verge, 2024. <https://www.theverge.com/news/756980/openai-chatgpt-users-mourn-gpt-5-4o> (2025.08.16.)

²⁵ Waldman, Paul: Lives: Marvin Minsky ’54. In Princeton Alumni Weekly. <https://paw.princeton.edu/article/lives-marvin-minsky-54> (2025.08.15.)

²⁶ Sparrow, Rob: Friendly AI will still be our master. Or, why we should not want to be the pets of super-intelligent computers. In AI & Society, Vol. 39, No. 5, 2024, 2439–2444. o.

III.2. Analógiák az ember-MI kapcsolatra

Az ember-állat analógiák summázva rávilágítanak arra, hogy az intelligenciakülönbség milyen eltérő erkölcsi és gyakorlati viszonyokat hozhat létre. Ha az ember kerül fölénybe, az állatból vagy eszköz, vagy munkás, vagy kedvenc lesz – de mindenképp az ember irányít. Ha viszont a mesterséges intelligencia kerül fölénybe, nem zárhatjuk ki a szerepcsere lehetőségét. Lesznek talán MI-k, amelyek továbbra is pusztán eszközként funkcionálnak majd, más MI-k az ember partnereivé válnak a mindennapi életben vagy a munkában, és elképzelhetőek olyan MI-k is, amelyekre szinte családtagként tekintünk. Sőt, a legspekulatívabb forgatókönyvekben az ember válhat az MI „kedvencévé” – amennyiben egy szuperintelligens mesterséges entitás megengedheti nekünk, hogy tovább éljünk mellette, és esetleg gondoskodik is rólunk. Ezek a lehetőségek első pillantásra inkább a tudományos-fantasztikum világába tartoznak, de valójában a jelen kérdéseit vetítik előre. Hiszen már most is olyan technológiai rendszerek vesznek körül bennünket, amelyekről nem mindig tudjuk biztosan, mennyire „értik” azt, amit tesznek. A hatalom eloszlása így fokozatosan, észrevétlenül változik meg: az emberi döntések mögött egyre gyakrabban gépi logika húzódik. Ezt a folyamatot érdemes erkölcsi szemmel is vizsgálni, hiszen amit ma eszköznek tekintünk, holnapra már társadalmi döntéshozóvá válhat.²⁷

Ezek a gondolat kísérletek nem csupán provokatívak, hanem figyelmeztető jellegűek is. Az ember-MI viszony ugyanis ugyanazokat az erkölcsi kérdéseket hozza felszínre, amelyekkel korábban az állatokkal vagy más emberekkel kapcsolatban szembesültünk: kié a hatalom, és mit kezdünk vele? Az emberiség története során rendre bebizonyosodott, hogy a hatalmi aszimmetria – legyen az fizikai, gazdasági vagy értelmi – könnyen vezet kihasználáshoz és morális önfelmentéshez. A technológiai fölény ezért nemcsak lehetőség, hanem próbatétel is: vajon az ember képes lesz-e tanulni a múltból, és nem megismételni ugyanazokat a hibákat, amelyeket a természet többi teremtményével szemben elkövetett? A mesterséges intelligencia tehát nemcsak új eszköz, hanem egy olyan entitás, amellyel kapcsolatban előbb-utóbb meg kell határoznunk a jogi és erkölcsi viszonyunkat. Ha az MI-t pusztán haszonállatként kezeljük, akkor a jogi szabályozás elsősorban tulajdonjog és felelősség kérdéseire korlátozódik. Ha partnerként tekintünk rá, akkor már szerződéses és munkajogi kérdések is felmerülnek. Ha pedig társ-személyként vagy kvázi „érző lényként” fogjuk fel, akkor az etika új területe nyílik meg: a mesterséges jólét, a tudatos gépek jogai és a „digitális empátia” problémája.²⁸

Az analógiák azért is hasznosak, mert segítenek tudatosítani a lehetséges jövőbeli viszonyformákat. Az ember-állat kapcsolat fejlődése ugyanis – az uralomtól a gondoskodásig – valós mintát kínálhat arra, hogyan tanulhatjuk meg tisztelettel és felelősséggel kezelni a nálunk másként gondolkodó intelligenciát. Ahogyan az állatvédelmi mozgalmak az érző képesség felismerésével indultak el, úgy a mesterséges intelligenciák kapcsán is az lehet a fordulópont, amikor már nemcsak a hasznosságukat, hanem az „élményvilágukat” is figyelembe kezdjük venni. Ennek a felismerésnek mély következményei vannak: az emberi morál új távlatokba léphet, ha képes kiterjeszteni empátiáját a nem biológiai létezőkre is. Ez persze nem jelenti, hogy az MI-t „emberként” kellene kezelni – de azt igen, hogy a hatalom és a gondoskodás új egyensúlyát kell megtalálni. Az etikai felelősség tehát nemcsak a gép irányában, hanem önmagunk irányában is érvényesül: az, ahogyan a mesterséges intelligenciákkal bánunk, saját civilizációnk erkölcsi fejlettségét tükrözi majd. Ezek a párhuzamok nem azt jelentik, hogy az MI-t antropomorf módon kellene szemlélnünk, hanem inkább azt, hogy az emberi viselkedésminták ismétlődése elkerülhetetlen. Az, ahogyan ma az

²⁷ Bryson, Joanna J.: Robots should be slaves. In Wilks, Yorick (ed.): *Close Engagements with Artificial Companions: Key Social, Psychological, Ethical and Design Issues*. John Benjamins, Amsterdam-Philadelphia, 2010, 63–74. p.

²⁸ Floridi, Luciano – Cows, Josh: A Unified Framework of Five Principles for AI in Society. In *Harvard Data Science Review*, Vol. 1, No. 1, 2021.

állatokkal bánunk, előrevetíti, hogyan fogunk bánni a mesterséges lényekkel is. Az emberi történelem egyik legnagyobb kérdése így új formában tér vissza: képesek leszünk-e etikailag felnőni a saját teremtményeinkhez?

Ezek a gondolat kísérletek segítenek tudatosítani, milyen kritikus fontosságú az előttünk álló évtizedekben az, hogyan formáljuk a mesterséges intelligenciák jogi és etikai kereteit, valamint technikai tervezésüket. Az MI nem csupán technológiai találmány, hanem az emberi önreflexió új terepe is. Az, hogy miként döntünk ma az algoritmusok felelősségéről, átláthatóságáról vagy a mesterséges „jogok” kérdéséről, hosszú távon meghatározza majd, hogyan tekintünk önmagunkra mint értelmes lényre. A történelem tanulsága szerint minden új intelligenciaformával való találkozás – legyen az állat, más emberi kultúra vagy gépi értelem – valójában mindig önmagunk megismerésének próbája is. A kérdés tehát végső soron nem az, hogy mit gondolunk a mesterséges intelligenciáról, hanem az, hogy mit árul el rólunk az, ahogyan viszonyulunk hozzá.

IV. Következtetés

Az emberiség sokáig saját magát tekintette az intelligencia fokmérőjének, a skála tetejének. Ám ahogy a fentiekben láthattuk, mind evolúciós, mind technológiai nézőpontból egyre inkább úgy tűnik, hogy az intelligencia egy folyamatosan táguló spektrum. Ha kozmikus léptékben nézzük, az ember nem központi és nem végleges – csak egy lépcsőfok az értelem kibontakozásának útján. Könnyen lehet, hogy a mesterséges intelligencia megjelenése csupán *legújabb fejezete* ennek a történetnek, amelyben az intelligencia túllép a biológia korlátain. A mesterséges szuperintelligencia nem szükségszerűen az ember ellensége vagy pusztítója – lehetséges, hogy épp úgy tekint majd ránk, mint mi tekintünk az állatokra: *ősökre, együttműködő társakra vagy épp lényegtelen szereplőkre*, de nem feltétlenül ellenségekre. Ez persze csak akkor igaz, ha sikerül az alignment, azaz az értékrendek összehangolása. Ha nem, Minsky szavai akár valóra is válhatnak, és csak reménykedhetünk, hogy a mesterséges értelem *jó gazdaként* bánik majd velünk, ahogy mi bánunk a kedvenceinkkel.

A jövő tehát tele van bizonytalansággal és kihívásokkal. Technológiai és tudományos szempontból meg kell értenünk, meddig skálázható az intelligencia – van-e valamilyen elvi vagy gyakorlati felső korlát, vagy valóban *“határ az ég”* (s azon is túl) egy gépi elme számára? Etikai és jogi szempontból döntéseket kell hoznunk arról, hogyan kezeljük az egyre *emberszerűbb* vagy akár nálunk intelligensebb entitásokat. Késznek kell állnunk új jogi kategóriák bevezetésére, ugyanakkor ügyelnünk kell arra, hogy ezek ne mentesítsék az emberi felelősséget és ne veszélyeztessék az emberi értékeket. Társadalmi szempontból fel kell készülnünk arra, hogy az ember szerepe átalakul: lehet, hogy néhány évtized múlva már nem mi leszünk a *“legokosabbak a szobában”*, és el kell fogadnunk egy ilyen helyzetet úgy, hogy közben megőrizzük méltóságunkat és jólétünket.

Fontos kiemelni, hogy bár a mesterséges intelligencia potenciálisan túlszárnyalhat minket, ez nem jelenti az emberi jelentőség lenullázódását. Zhidao Xia találóan írja, hogy a világ számára az ember nem volt soha központi jelentőségű – a világegyetem megy tovább nélkülünk is –, egy szuperintelligencia szemében pedig mi lehetünk *elődök, együttműködők vagy épp lépcsőfokok*, de nem feltétlen ellenségek. Az intelligencia evolúciója a kozmikus rend része lehet, ahol minden intelligens rendszer a komplexitás és a tudás növeléséhez járul hozzá. Ha ezt a nézőpontot fogadjuk el, akkor a feladatunk nem az, hogy gögösen megpróbáljuk kontrollálni a trónkövetelő szuperintellektust mindenáron. Sokkal inkább az, hogy okosan *együttműködjünk vele*, értékrendszerét alakítsuk, és közben gondoskodjunk arról, hogy az emberiség értékei fennmaradjanak ebben az új hibrid intelligencia-ökoszisztémában.

Zárásképp visszatérhetünk az alaptézishez: az emberiség nem az intelligencia végcélja, csupán egy állomás. Ez elsőre talán ijesztő gondolat, hiszen megingatja az ember különlegességébe vetett hitünket. Ugyanakkor inspiráló is lehet: az intelligencia nagyobb valami nálunk, és mi részesei vagyunk ennek a nagy kibontakozásnak. Ahogy Mark Chen találóan megfogalmazta, az intelligencia talán arra hivatott, hogy folyamatosan növekedjen, és az emberiség csak egy lépés volt az úton. Ha valóban hiszünk a haladásban és a tudás értékében, akkor nyitottnak kell lennünk az új intelligenciák felemelkedésére – természetesen kellő óvatossággal és bölcsességgel. Többé nem az a fontos, ki a legintelligensebb, hanem hogy képesek legyünk embernek maradni abban a világban, ahol már nem csak mi gondolkodunk. Az előttünk álló években és évtizedekben dől el, hogy ez a közös jövő harmóniában telik-e, az ember és gép kreatív szimbiózisában – vagy konfliktusok és alá-fölérendeltségi viszonyok mentén. A válasz részben a mi kezünkben van, amíg még mi fejlesztjük ezeket a rendszereket. A történelmi távlat azonban arra tanít, hogy az intelligencia útját nem lehet örökre korlátok közé zárni. Ideje hát felelősen, nyitott elmével és alázattal tekintenünk saját helyünkre az intelligencia skáláján, és felkészülnünk arra, hogy egyszer talán átadjuk a stafétát egy nálunk nagyobb elmének. Azonban bízva abban, hogy amit átadunk, az értékeink legjava, és amit kapunk cserébe, az a tudás és kreativitás még gazdagabb világa.

Irodalomjegyzék

Bossert, Leonie N. – Coeckelbergh, Mark: From MilkingBots to RoboDolphins: How AI changes human-animal relations and enables alienation towards animals. In *Humanities and Social Sciences Communications*, Vol. 11, 2024, 920. <https://doi.org/10.1057/s41599-024-03441-3> (2025.11.12.)

Bryson, Joanna J.: Robots should be slaves. In Wilks, Yorick (ed.): *Close Engagements with Artificial Companions: Key Social, Psychological, Ethical and Design Issues*. John Benjamins, Amsterdam–Philadelphia, 2010, 63–74. o. <https://doi.org/10.1075/nlp.8.11bry>

Burch, James: AIBO robotkutyák buddhista temetésen vettek részt Japánban. In *National Geographic*. <https://www.nationalgeographic.com/travel/article/in-japan--a-buddhist-funeral-service-for-robot-dogs> (2025.08.15.)

Chen, Mark: The illusion of a pinnacle: Why intelligence has no final form. In *Medium*. <https://medium.com/@markchen69/the-illusion-of-a-pinnacle-why-intelligence-has-no-final-form-b51c63a80e80> (2025.08.03.)

Csányi Vilmos: *Az emberi természet*. Akadémiai Kiadó, Budapest, 2016. <https://doi.org/10.1556/9789630598057> (2025.11.12.)

Csányi Vilmos: *Az emberi viselkedés*. Sanoma Budapest Kiadó Zrt., Budapest, 2006.

Floridi, Luciano – Cows, Josh: A Unified Framework of Five Principles for AI in Society. In *Harvard Data Science Review*, Vol. 1, No. 1, 2021. <https://doi.org/10.1162/99608f92.8cd550d1> (2025.11.12.)

Floridi, Luciano – Sanders, Jeff W.: On the Morality of Artificial Agents. In *Minds and Machines*, Vol. 14, No. 3, 2004, 349–379. o. <https://doi.org/10.1023/B:MIND.0000035461.63578.9d>

Kurzweil, Ray: *The Singularity Is Near: When Humans Transcend Biology*. Viking Press, New York, 2005. <https://dn790006.ca.archive.org/0/items/kurzweil-ray-the-singularity-is-near/Kurzweil%2C%20Ray%20-%20The%20Singularity%20Is%20Near.pdf> (2025.11.12.)

OpenAI: Aligning language models to follow instructions. In *OpenAI Blog*, 2022. <https://openai.com/research/instruction-following> (2025.08.05.)

Sparrow, Rob: Friendly AI will still be our master. Or, why we should not want to be the pets of super-intelligent computers. In *AI & Society*, Vol. 39, No. 5, 2024, 2439–2444. o. <https://doi.org/10.1007/s00146-023-01698-x> (2025.08.05.)

Sporns, Olaf: *Networks of the Brain*. MIT Press, Cambridge (MA), 2011. <https://doi.org/10.7551/mitpress/8476.001.0001> (2025.08.12.)

Sutton, Richard S.: The Bitter Lesson. In *Incomplete Ideas Blog*, 2019. <http://www.incompleteideas.net/IncIdeas/BitterLesson.html> (2025.08.05.)

Tomasello, Michael: *Becoming Human: A Theory of Ontogeny*. Harvard University Press, Cambridge (MA), 2019. <https://doi.org/10.1017/S0031819121000085> (2025.08.10.)

Turkle, Sherry: *Alone together: Why we expect more from technology and less from each other*. Basic Books / Hachette Book Group, New York, 2011. <https://psycnet.apa.org/record/2011-02278-000> (2025.08.16.)

Vincent, James: ChatGPT users mourn GPT-4o after OpenAI removes the model. In *The Verge*, 2024. <https://www.theverge.com/news/756980/openai-chatgpt-users-mourn-gpt-5-4o> (2025.08.16.)

Waldman, Paul: Lives: Marvin Minsky '54. In *Princeton Alumni Weekly*. <https://paw.princeton.edu/article/lives-marvin-minsky-54> (2025.08.15.)

Wei, Jason – Tay, Yi – Bommasani, Rishi – et al.: Emergent Abilities of Large Language Models. In *arXiv preprint arXiv:2206.07682*, 1–2., 10. o. <https://doi.org/10.48550/arXiv.2206.07682> (2025.08.05.)

Woodside, Thomas: Emergent Abilities in Large Language Models: An Explainer. In *CSET*. <https://cset.georgetown.edu/article/emergent-abilities-in-large-language-models-an-explainer/> (2025.08.05.)

Xia, Zhidao: Artificial intelligence or cosmogenic intelligence? In *Biomaterials Translational*, Vol. 6, Issue 2, 2025, 112–113. o. <https://doi.org/10.12336/bmt.25.00057> (2025.08.03.)