

Mesterséges intelligencia

INTERDISZCIPLINÁRIS E-FOLYÓIRAT

OPEN ACCESS



DOI <https://www.doi.org/10.35406/MI.2025.2.1>

ISSN 2676-9611

VII. évfolyam 2025/2. szám

WEB: www.kpluszf.com

K+F STÚDIÓ Kft.

az



IMPRESSZUM

MESTERSÉGES INTELLIGENCIA

Interdiszciplináris e-folyóirat

Alapítva: 2019-ben.

ISSN 2676-9611

A Nemzeti Média- és Hírközlési Hatóság Hivatala a médiaszolgáltatásokról és a tömegkommunikációról szóló 2010. évi CLXXXV. törvény 46.§ (4) bekezdése alapján nyilvántartásba vett sajtótermék (határozatról szóló értesítés iktatószáma: CE/5420-5/2019).

A *Mesterséges intelligencia* interdiszciplináris e-folyóirat a K+F Stúdió Kft. által, társadalmi felelősségvállalási (CSR) stratégia keretében alapított és kiadott, negyedévente megjelenő Open Access (nyílt hozzáférésű) internetes periodika, melyben két anonim és két nem anonim szakmai lektor bírál minden tanulmányt.

A Kiadó adatai:

Kiadó: K+F Stúdió Kft.

A kiadó székhelye: 4032 Debrecen, Tarján utca 55.

Mobil: +36-30-4849779

E-mail: info@kpluszf.com

Web: www.kpluszf.com

Kiadásért felelős személy: Mező Katalin (PhD)

A Szerkesztőség adatai:

Levélcím: K+F Stúdió Kft., 4032 Debrecen, Tarján utca 55.

Mobil: +36-30-4849779

E-mail: info@kpluszf.com

Web: www.kpluszf.com

Alapító főszerkesztő: Mező Ferenc (PhD)

Tördelő szerkesztő: Mező Katalin (PhD)

Együttműködő civil szervezet:

Kocka Kör Tehetséggondozó Kulturális Egyesület (www.kockakor.hu)

Szerkesztőség (ABC rendben):

Bodnár Gabriella, (PhD, habil., Soproni Egyetem)

Gyarmati Péter (Prof. Dr.)

Kelemen Lajos (PhD, OKOSKOCKA Kft.)

Mező Ferenc (PhD, K+F Stúdió Kft.)

Mező Katalin (PhD, Debreceni Egyetem)

Orbán Réka (PhD, Babes-Bolyai Egyetem)

Pénzes Dávid (PhD, Káldor Miklós Kollégium)

Pšenáková Ildikó (PhD, Trnava University in Trnava, Szlovákia)

Pšenák Péter (Ph.D., Comenius University Bratislava, Szlovákia)

Simó Ferenc Zoltán (dr., LL.M, Debreceni Egyetem)

Szabóné Balogh Ágota (PhD, Gál Ferenc Főiskola)

Szűts Zoltán (PhD, Eszterházy Károly Katolikus Egyetem)

Tomac, Zvonimir (PhD, University J.J. Strossmayera of Osijek, Horvátország)

Vass Vilmos (PhD, habil., Budapesti Metropolitan Egyetem, Selye János Egyetem)

Külön nem hivatkozott illusztrációk forrása: <https://pixabay.com>

TARTALOM

SZERKESZTŐI KÖSZÖNTŐ.....	5
ELMÉLETI ÉS EMPIRIKUS TANULMÁNYOK	
THEORETICAL AND EMPIRICAL STUDIES	7
Gulyás, Géza and Petrovánszki, Julianna: THE SAME SYSTEM WITH DIFFERENT AND ADDITIONAL SUPPORT SYSTEMS – THE POSSIBLE FUNCTIONS OF AI IN WARGAMING	9
Oravecz, Adrienn: A BIZTONSÁG MEGÍTÉLÉSÉNEK VIZSGÁLATA MOZGÁSSÉRÜLT VEZETŐK ESETÉBEN – A MESTERSÉGES INTELLIGENCIA SZEREPÉNEK KITERJESZTÉSÉVEL.....	35
Sultan, Darwish; Houssein, Zairi and Kocsis, Dénes: DATA-DRIVEN TENSOR PRODUCT POLYTYPIC BASED MODEL FOR ENVIRONMENTAL NOISE MONITORING	45
Gyarmati G. Péter: GONDOLATOK A MESTERSÉGES INTELLIGENCIÁRÓL TANÁROKNAK	57
Csajbók Gábor: MESTERSÉGES INTELLIGENCIA ÉS KREATIVITÁS: HATÁROK, KÜLÖNBSÉGEK ÉS ÚJ FOGALMI KERETEK	65
MÓDSZERTANI TANULMÁNYOK	
METHODOICAL STUDIES.....	75
Horváth, Dávid: WHEN ARTIFICIAL INTELLIGENCE GENERATED CONTENT BECOMES WEAPON-GRADE COMMUNICATION	77
Szabó Tibor, Hegedűs Orsolya és Németh Zsófia: KREATÍV TÖRTÉNETMESÉLÉS OZOBOTOK SEGÍTSÉGÉVEL	99
Papp Sándor és Mező Kristóf Szíriusz: GITÁRZENE A MESTERSÉGES INTELLIGENCIA KORÁBAN – GONDOLATOK A XV. GYÖNGYÖSI NEMZETKÖZI GITÁRFESZTIVÁL KAPCSÁN	109

Muntean, Loredana and Koren, Morel: INTEGRATING DIGITAL PEDAGOGY INTO TRADITIONAL MUSIC EDUCATION	117
--	-----

MŰHELY, RENDEZVÉNY WORKSHOPS AND EVENTS.....	127
---	------------

Szűts Zoltán és Szőke-Milinte Enikő: MI ÉS AZ NTP-TEHETSÉG-25-0142 PROJEKT	129
---	-----

Mező, Ferenc: SHORT REPORT ABOUT THE 'LEARNING AND SOCIETY (2025)' CONFERENCE AND THE MEC-SZ-24-149026 PROJECT	135
--	-----

A GÁBOR DÉNES EGYETEM MIT ÉS A STRT HOLDING INGYENES ONLINE MI TANANYAGA	139
---	-----

FELHÍVÁS INTERDISZCIPLINÁRIS JUNIOR KUTATÓCSOPORTBA TÖRTÉNŐ BEKAPCSOLÓDÁSRA	143
--	-----

CODE POETRY PÁLYÁZAT (2026).....	147
----------------------------------	-----

SZERKESZTŐI KÖSZÖNTŐ



Tisztelt Olvasó!

Üdvözljük a *Mesterséges intelligencia* folyóirat VII. évfolyam 2025/2. számának megjelenése alkalmából!*

A korábbiakhoz hasonlóan a mesterséges intelligencia sokszempontú megközelítésével találkozhatunk jelen lapszámunkban is.

Az első tanulmányban Gulás Géza ezredes és Petrovánszki Julianna a mesterséges Intelligencia szereplehetőségeit mutatják be a hadi játékok terén.

Oraveczi Adrienn a mozgássérült sofőrök aspektusából a mesterséges intelligencia közlekedésbiztonsággal kapcsolatos szerepére irányítja a figyelmet.

Darwish Sultan, Zairi Houssem és Kocsis Dénes okosváros-irányítópultokba adaptálható, a valós idejű környezeti zajkezeléshez, ezáltal az adaptív várostervezéshez hasznosítható MATLAB-alapú megoldást mutat be.

***Az MI témakörrel ismerkedők számára bevezető tanulmányként javasoljuk: Mező Ferenc és Mező Katalin (2019): Interdiszciplináris kapcsolódási lehetőségek a mesterséges intelligenciára irányuló cél-, eszköz- és hatásorientált kutatáshoz. *Mesterséges intelligencia – interdiszciplináris folyóirat*, I. évf. 2019/1. szám. 9–29. Doi:**

<https://www.doi.org/10.35406/MI.2019.1.9>

Ezt követően Gyarmati Péter tanároknak szóló, mesterséges intelligenciával kapcsolatos gondolataiba nyerhetünk bepillantást.

Az elméleti tanulmányok körét Csajbók Gábor tanulmánya zárja, melyben a mesterséges intelligencia és kreativitás témakörben tesz továbbgondolásra, kutatásra alkalmas felvetéseket.

A módszertani tanulmányokat közt elsőként Horváth Dávid a mesterséges intelligencia által generált fegyvernek minősülő kommunikáció témakörébe enged bepillantást.

Szabó Tibor, Hegedűs Orsolya és Németh Zsófia edukációs robotok kreatív történetmesélésben (s ezáltal lényegében bármilyen tantárgyi tartalom játékosításába, élményszerűbbé tételébe) történő felhasználási lehetőségére ad módszertani példát.

A harmadik és negyedik módszertani tanulmány a mesterséges intelligencia és a zene témakörök keresztszűrésére fókuszálnak. Papp Sándor és Mező Kristóf Szíriusz a gitárzene, Loredana Muntean és Morel Koren a hagyományos zeneoktatás kontextusában közelíti meg a mesterséges intelligencia témát.

A műhelyeket, rendezvényeket bemutató rovatban elsőként Szűts Zoltán és Szőke-Milinte Enikő mutatja be a mesterségei intelligencia és az NTP-TEHETSÉG-25-0142 projekt közötti kapcsolatot.

Ezt követően a Mező Ferenc a MEC-SZ-24-149026 projekt keretében megvalósult Ta-

nulás és Társadalom (2025) konferenciáról ad közre összefoglalót.

Ebben a rovatban található továbbá a Gábor Dénes Egyetem Mesterséges Intelligencia Tudásközpont és a STRT Holding ingyenes MI tananyagára vonatkozó felhívás is.

A műhely, rendezvény bemutatását célzó rovatot két felhívás zárja. Az egyik Interdiszciplináris Junior Kutatócsoportba történő bekapcsolódásra biztatja az Olvasókat, a második pedig a Code Poetry (2026) pályázati felhívást tartalmazza.

Gondolatébresztő és tanulmány beküldésére motiváló olvasást kíván Önnek a Szerkesztőség nevében is:

*Dr. Mező Ferenc
alapító főszerkesztő*

ELMÉLETI ÉS EMPIRIKUS TANULMÁNYOK

THEORETICAL AND EMPIRICAL STUDIES

**THE SAME SYSTEM WITH DIFFERENT AND ADDITIONAL SUPPORT
SYSTEMS – THE POSSIBLE FUNCTIONS OF AI IN WARGAMING**

Author(s) / Szerző(k):

Géza Gulyás (Ph.D., Col.)

MoD Defence Innovation and Capability Development Department;
University of Public Service
(Hungary)

Julianna Petrovánszki

MoD Defence Innovation and Capability Development Department
(Hungary)

E-mail:

petrovanszkijuli@gmail.com

Cite: Gulyás, Géza and Petrovánszki, Julianna (2025): The same system
Idézés: with different and additional support systems – The possible
functions of AI in wargaming. *Mesterséges Intelligencia – interdiszciplináris
folyóirat*, VII. évf. 2025/2. szám. 9-33.
Doi: <https://www.doi.org/10.35406/MI.2025.2.9>



This work is licensed under a Creative Commons Attribution-
NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0007

Reviewers: Public Reviewers / Nyilvános Lektorok:
Lektorok: 1. Kaiser Ferenc (Ph.D.), Nemzeti Közszerológati Egyetem
2. Harangi-Tóth Zoltán (Ph.D.), Nemzeti Közszerológati Egyetem

Anonymous reviewers / Anonim lektorok:
3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

UGYANAZ A RENDSZER KÜLÖNBÖZŐ ÉS KIEGÉSZÍTŐ TÁMOGATÓ RENDSZEREKKEL – AZ MI LEHETSÉGES FUNKCIÓI A HADIJÁTÉKOKBAN

A professzionális hadijátékok egyfajta reneszánszukat élik, mint elismert és rendszeresített eljárás mind hazai, mint NATO-s szinten. Ezzel párhuzamosan az egyre inkább teret hódító mesterséges intelligencia (MI) is kiemelt kérdéskörre váltak a védelmi szférában. Jelen tanulmány a szakirodalmat áttekintve azon kérdésre keresi a választ, hogy MI-alapú eszközök és rendszerek miként és hogyan alkalmazhatóak a professzionális hadijátékozás területén, különös tekintettel a potenciális veszélyekre és korlátokra. MI-rendszerek elősegítik és gyorsítják a hadijátékok tervezését, lejátszását, illetve annak elemzését is, ugyanakkor a szerzők hangsúlyozzák, hogy ezt csak bizonyos korlátozásokkal érdemes alkalmazni. A szerzők rámutatnak, hogy háború, és így a hadijátékozás emberi tényezője továbbra is meghatározó marad, a mesterséges intelligencia nem helyettesíti, hanem helyesen alkalmazva támogatja azt.

Kulcsszavak: mesterséges intelligencia, hadijáték, döntéshozatal

Diszciplína: hadtudomány, informatika

Abstract

Professional wargames are undergoing a renaissance as a recognized and systematized procedure at both the domestic and NATO levels. At the same time, artificial intelligence (AI), which is becoming increasingly prevalent, has also become a key issue in the defense sector. This study reviews the literature to find answers to the question of how AI-based tools and systems can be applied in the field of professional wargames, with particular regard to potential risks and limitations. AI systems facilitate and accelerate the planning, design, execution, and analysis of military wargames, but the authors emphasize that their use should be subject to certain restrictions. The authors point out that war, and thus the human factor in wargames, remains decisive; artificial intelligence does not replace it, but rather supports it when used correctly.

Keywords: artificial intelligence, wargaming, decision making

Discipline: military science, informatic

In this short article, the authors discuss the relationship between wargames and artificial intelligence (hereinafter referred to as AI) and try to figure out if it is worth it or even right

to use AI in wargames. It may seem surprising at first, but this question is not primarily a technical one, but rather a civilizational one. Why? Because the application of AI in war-

games is actually just a testing ground for how far humans are willing to outsource their thinking and the responsibility for their own thinking. Warfare has always been the most acute form of human decision-making; if the human mind begins to disappear there, sooner or later it will follow suit in all other areas...

Before delving into the questions surrounding AI, it might be useful to spend a few paragraphs explaining wargaming itself. These terms are so in vogue today in military science, lectures, and articles. They are mentioned and emphasized in countless places, but do we know and understand what they actually mean? Well, this is not such a simple question, because if, for example, we search for Hungarian-language articles and studies on the internet, we will find that there are few such writings and authors. Unfortunately, this means that it is not nearly as common a topic in Hungarian military culture as it is in other armed forces. At the same time, the alliance is placing increasing emphasis on the use of wargames to help maintain its deterrence and resilience against potential adversaries.

With regard to AI the NATO 2022 Strategic Concept states that emerging and disruptive technologies (or, more simply: emerging technologies), including AI, present both opportunities and risks that affect the nature of conflicts and have strategic significance (NATO, 2022).

If we look up the term wargaming or wargame, we find countless definitions. Perhaps one of the most comprehensive and easily understandable definitions for “non-experts” can be found in a synonym

dictionary. According to this: “Wargaming is an activity in which military forces or other organizations participate in a simulated war or conflict situation in order to test their strategies and decision-making processes. During this process, participants practice military operations in a virtual or real environment, taking into account various factors such as terrain, weather, enemy strength, and their own resources...” (Szinonimák.hu). So, this is an activity performed by soldiers. Basically, yes, but the issue is not that simple. A few lines later, the editors emphasize that “wargames are not only useful for military forces, but also for other organizations and institutions...” (Szinonimák.hu). Indeed, it can be said that the principles and procedures developed over the centuries can be useful for anyone or any organization that wants to “foresee” future processes, decisions, or their consequences as accurately as possible. However, the authors base their approach on the conceptual approach set out in the NATO Handbook, according to which: “Wargames are representations of conflict or competition in a safe-to-fail environment, in which people make decisions and respond to the consequences of those decisions” (NATO, 2023, p.8.). As can be seen from this, human input is a central element of wargaming, which consists of the following:

- the players,
- the decisions they make,
- the narratives they create and shape,
- the experiences they share,
- the lessons they learn (Doktrinzentrum der Bundeswehr, 2024).

The “safe-to-fail” environment characteristic of wargames refers to the fact that the decisions made by people have no direct physical impact on reality, meaning that it is possible to explore, within a safe framework, where the various decision alternatives lead, what their effects and risks are, and how and why a given situation can be lost or won, which are the characteristics that justify the emerging trend of wargames (UK Ministry of Defence, 2017).

In the following chapters of this study, after a brief historical overview, we will examine whether – taking into account the risks and limitations – artificial intelligence can be used in wargaming as it is currently practiced, and if so, in which areas. In addition to outlining the possibilities, we will, of course, pay special attention to the potential negatives, dangers, and ethical issues, and we hope that in the last chapter, the reader will find logical answers to the questions posed in the first paragraph.

The early days of wargaming

Previously, we referred to a process spanning several centuries. The question may arise: where did this “science” begin? This is difficult to determine, since – quite understandably – since the dawn of human history, all warring parties have desired and continue to desire to know the plans and future activities of their opponents, so that they can take appropriate and effective countermeasures against the actions of enemy... but above all, they have wanted to surprise the enemy with moves that would

allow them to seize and maintain the initiative. In almost all cases, this is the key to success and victory.

Several board games known today date back to ancient times, the most famous of which is chess, which remains popular to this day. In some respects, these can be considered the precursors of wargames, but since they were all quite abstract, we cannot really speak of wargames in the modern, professional sense (Somkuti, 2021). Among military themes, one of the earliest examples was the Indian game Chaturanga, which dates back more than two thousand years and included the Indian weapons of the time: elephants, cavalry, chariots, and infantry (Sabin, 2014). Jumping forward in time, we arrive at the early 1800s... to Napoleon, who was probably influenced by a board game that appeared a few decades earlier, which, albeit in a rudimentary form, featured various factors such as terrain, weather, logistics, and dice.

According to records, the French ruler was one of the first to use different coloured blocks to designate troops, and since 1811, alongside Georg von Rösswitz’s numerous scientific innovations, blue has been used to designate friendly forces and red to designate the enemy. Rösswitz’s other innovation was that for instance, unlike in board games, teams did not immediately fall off the board but remained on the “battlefield” until their destruction, albeit with continuously decreasing strength or numerical value. Subsequently, after the Franco-Prussian War ended in a sweeping Prussian victory, the procedure spread among the world’s armed

forces, and the emergence of professional wargames for military use in the modern sense can be traced back to the Prussian Kriegsspiel of 1824. Earlier games related to the theme of war were primarily designed for entertainment purposes, while Kriegsspiel, which focused on military decision-making and was based on topographical maps, was simulative in nature and was able to convey the fog of war (Nebels des Krieges) to participants (Pielström and Wintjes, 2019).

Prior to World War I, the future adversaries conducted numerous wargames simulating potential wars. For instance, even before the world war, wargames were part of the planning of major operations based on German doctrines (Caffrey, 2019). During the Great War, the interwar period, and World War II, wargames were used as a respected tool. The US Navy took this to such an extent during World War II that, according to them, nothing came as a surprise to them during the war except for the kamikaze pilots (UK Ministry of Defence, 2017).

After World War II, wargames underwent a transformation. In Western countries, led by the US, interest in wargames declined and shifted toward computerization. In the 1960s, wargaming regained popularity and took on a political-military dimension. In the 1970s and 1980s, it continued to develop and become more institutionalized, with a focus on computer-assisted analysis coming to the fore (Caffrey, 2019). Simultaneously, civilian wargames created for entertainment purposes also enjoyed great popularity starting in the 1970s.

At that time, on the eastern side of the Iron Curtain, wargaming continued to develop less in the direction of computerization and more in the traditional direction. In the Soviet Union, data collected from World War II was used to increase realism and make wargaming as realistic as possible (Caffrey, 2019). Hungarian wargaming, like military decision-making as a whole, was characterized by Soviet procedures, which placed emphasis on well-prepared military decision-makers, while the relatively small staff simply relayed orders to subordinate units. In this process, wargames were intended to develop the decision-making abilities of military leaders on the one hand, while on the other hand preparing the staff to implement the above procedures (Harangi Tóth, 2019). In the final years of the Cold War, in a spirit of *détente*, joint US-Soviet military exercises were held, in which wargames were used as a confidence-building measures (Caffrey, 2019).

Currently, the use of wargames is highly regarded at NATO level, and at the same time, an increasing number of member states are implementing capability developments in this area. In the Hungarian Defence Forces, wargames are used as part of the Military Decision Making Process, mainly in the comparison of developed courses of action (Harangi Tóth, 2019). It is also part of Hungarian military higher education, where it appears from BSc to the Military Senior Leadership Course.

In terms of their use, we can distinguish between professional and commercially available wargames. However, due to the focus of

this article, the authors present the phenomenon of wargaming without claiming to be exhaustively comprehensive. Within professional wargames, we primarily refer to those used by the defence sector, but wargames are also present in business and other policy areas. Military applications by the defence sector can be categorized based on three fundamental characteristics: purpose of use, level of military operations, and time period appearing in the scenario (Caffrey, 2019). Depending on the purpose of use, they can be analytical or educational. The purpose of analytical wargames is to support decision-making, typically by facilitating the systematic analysis of complex situations, revealing new connections, interactions, and possible consequences, or by generating data. Educational wargames help develop decision-making skills and are a useful tool for both current and future leaders.

But what can wargaming be used for in practice? The Handbook of the Bundeswehr, for instance, defines the following areas of application: decision-making support, increasing mental resilience, and innovation (see: Doktrinzentrum der Bundeswehr, 2024). In supporting decision-making, wargaming appears as a qualitative method, enabling comprehensive examination of complex problems and providing a better information base for decision-makers, for instance, the aforementioned “CoA wargaming.” This is particularly effective when combined with quantitative methodology. The use of wargaming increases mental resilience through cognitive learning. Factors such as fear of failure, poor

organizational culture, or social pressure have a negative impact on decision-making (see: Doktrinzentrum der Bundeswehr, 2024). Organizations operating under strict hierarchical structures, such as the armed forces, are characterized by a culture of victory, in other words, intolerance of defeat and mistakes, as well as an unconstructive attitude towards the questions and opinions of subordinates in the hierarchical system. In the field of innovation, wargames can promote the development of new ideas and approaches, as a well-organized and well-facilitated wargame can enable participants of different ranks and positions to work together and share their opinions, thereby strengthening creative thinking (Doktrinzentrum der Bundeswehr, 2024).

The potential roles of artificial intelligence in wargames

It is worth clarifying the place of AI in our world and some basic questions at the beginning of this study. Artificial intelligence is not a piece of software or a specific product, but rather a scientific discipline, an increasingly independent and growing field of science. For the sake of clarity, it is worth approaching the issue in this way. Another important question is the existence of AI itself.

Contrary to the general belief, we can currently only work with Artificial Narrow Intelligence (ANI) systems, which, when used properly, are quite effective, but are also task-

specific and cannot be applied to multiple tasks without human input or intervention. It is worth to note that although large language models (LLMs), which are frequently used today, appear to be generally applicable, they are in fact part of narrow AI. To our knowledge, artificial general intelligence (AGI) and artificial super intelligence (ASI) currently exist only in theory, and it is uncertain whether human science will be able to create AGI and subsequently ASI (USMA Library). Just as we can only have vague ideas at present about whether the answer is yes, we also have vague ideas about when and how this might happen. Despite the many questions, we must also consider these possibilities, so it is possible (and certainly worthwhile) to design a wargame with the specific aim of mapping out the potential effects and dangers of AGI or ASI with the involvement of experts. It may already be apparent from these few thoughts that the relationship between AI and wargames is two-way: we can use wargames to examine the development of AI systems and their impact on warfare, for instance, but we can also use AI tools to design, execute, and analyse wargames. It is important to note that the primary purpose of these systems is to supplement and support human input, not to replace it.

Lastly, it might be worth raising a slightly philosophical point, namely that in the future we will probably have to distinguish between strong and weak AI. Strong AI refers to truly intelligent, self-aware intelligence, while weak AI pretends to be intelligent (University of Helsinki and MinnaLearn, 2018). At present,

we can only talk about weak AI systems, but the future will definitely see the creation of “strong” systems.

Below, we will dive into how AI can be used in the various phases of wargames, without claiming to be exhaustive.

*Planning and designing wargames,
creating scenario*

The creation of different scenarios is a crucial issue, because scenario generation is the point, and the opportunity where AI truly enters the realm of wargame creation, not merely as an analyst, but as a story-creating partner. Humans can intuitively come up with one scenario, maybe two or three... but AI can generate a huge number of them in a short period of time. It does this by taking into account the enemy's previous behaviour, various reconnaissance and intelligence data, logistical and weather factors, and basically anything else.

The narrative of a wargame is essentially an imagined scenario, a potential conflict situation that players bring to life with their decisions. If this scenario is created by artificial intelligence, then it becomes the “scriptwriter” and, at the same time, the shaper of the framework for thinking at all three levels of warfare in a way that is more complex than human capabilities and, in fact, probably more accurate and much faster.

And now consider how AI should support existing military wargaming systems at three different levels of warfare.

At the strategic level, the task is actually to create an entire world, with a fictional or real-

life political background, federal relations, economic conditions, time frame, theatre of war, etc. Here, the role of AI is to collect data, simulate the consequences of political decisions, and create realistic geopolitical narratives. Nevertheless, current AI systems often over-simplify complex political processes, tending to underestimate the role of irrationality, brain-storming, and intuition in human decision-making... and this is certainly a serious limitation.

At the operational level, the composition and position of friendly and enemy forces, command intent, objectives, physical and temporal constraints of the theatre of operations, as well as logistical, communications, morale, weather, and other conditions are determined.

An AI method must assist the staff in the creation of simulated counterparts of real units, vehicles, and tools, weather and terrain modelling, automatic force deployment based on the characteristics of real weapon systems and doctrinal principles, as well as previously learned knowledge and force assessment. AI methods are already easy to use in these areas today, as the data is available and the models can be taught. The problem here is rather that AI does not always understand “intent,” meaning that it knows what can happen, but does not always understand why.

At the tactical level, AI may be capable of performing realtime tactical simulations (movement of units at the squad, platoon, or battalion level, fire support, etc.), even in space, and modeling probable outcomes (e.g., combat, supply shortages, detection inter-

ference). As a result, it must be able to redesign the scenario if the situation changes unexpectedly (e.g., logistical failure, weather disaster), thereby offering staff members alternatives that they might not have thought of themselves, for instance: “What happens if the enemy attacks six hours earlier?” “What happens if ammunition supplies are reduced by 20%?” etc.

As a result of using AI methods, military staff can see not just two or perhaps a few, but hundreds of possible scenarios – the worst-case scenario and the most likely scenario must always be worked out –, which they can play out in a matter of minutes, as a simulation within the whole procedure. This generates a huge amount of new data and information in almost every case, which also needs to be analysed and evaluated with the help of AI, for example, by searching for patterns that remain hidden from the human eye and mind.

Execution of wargames

Accelerated probability analysis. In traditional wargames, participants explore only a few courses of action, but AI can generate and play out many more in a matter of minutes. The Monte Carlo method is a stochastic simulation method that uses computer tools to generate the final result of a given experiment, after which the numerical characteristics obtained as a result are recorded and evaluated. The error in the result is determined by calculating the standard deviation. Using Monte Carlo simulation, it can generate probability distributions: e.g., the probability of success, failure, or partial

success. This gives decision-makers a more statistical-based picture of which alternative is the most viable, but for now, these are just numbers...

The use of AI models can be useful during a wargame, assisting the staff in assessing the probability of success of the actions proposed by the participants, thereby “lifting” some of the burden from the facilitators and adjudicators and making the scenario flow more smoothly.

One common reason for scepticism about wargames is the use of dice, which are often used to determine the outcome of a decision. However, it is important to know and understand that wargame designers do not use this method arbitrarily but rather based on mathematical probability calculations and modifying factors related to the given topic, while also incorporating the phenomenon of chance into the mechanics of the wargame. In general, war-games are designed so that numerous factors can modify the outcome of a dice roll, such as resource allocation, differences in capabilities, circumstances, and so on (see: Harangi Tóth, 2020). This is therefore the result of thorough work by wargame designers based on scientific knowledge, taking into account the aspects of realism and the functionality and playability of wargames... and by no means just “guesswork”!

Unfortunately, in today's seemingly digitized, sometimes “over-digitized” world – although the bitter reality of the Russo-Ukrainian war provides countless examples to the contrary – when we try to improve everything sometimes recklessly with various

computer systems, the mention of dice is tantamount to frivolity for many military leaders, as for many it may initially bring to mind board games designed for children, where the success of each decision seems to depend solely on luck (Petrovánszki, 2023). An AI-based tool could be a solution here too, as participants are more likely to accept the “result thrown” by artificial intelligence (since it is already “digitized” and therefore acceptable) than if we did the same thing with dice, due to their (sometimes excessive) trust in technology.

Western military culture, especially the American and British approaches, relied much more on intuition, situational awareness, and decentralized initiative, accepting that in reality not every decision can be optimized, while AI – like Soviet-style commanders – seeks to optimize, always looking for the best combination. The Soviet Union's decision-making culture was highly centralized and formalized: they tried to measure every variable. They were much better at this than their Western counterparts, and it was effective in the short term (e.g., artillery planning, logistical synchronization), but in the longer term, it made thinking rigid, unable to keep up with changes.

At present, AI does not seem to understand that human commanders/leaders often deliberately choose suboptimal solutions when they prioritize surprise or psychological superiority over rationalism. The biggest methodological dilemma of AI-supported wargames is: what do we measure and how, if the “best decision” is not always what the

algorithm considers optimal? It is also worth considering what “best decision” actually means. After all, while AI takes into account and examines loss minimization, the time factor (achieving the goal as quickly as possible), the efficiency of resource utilization, and the level of risk, human decision-making also involves unexpected events, psychological effects, moral and political issues, or the conscious desire to “not choose the most obvious option.”

It seems that this is the borderline... AI decisions may be rational, but they may not necessarily be wise, which is also a powerful warning! A “wise decision” may not necessarily yield the best result but rather attempts to envision the least bad future. This is a different type of thinking, as it does not optimize, but rather anticipates.

The correctness of decisions cannot be quantified or measured precisely, and we believe that it is not worth attempting to do so, even with the help of AI. Instead, the focus should be on examining the correctness of the decision-making process, on how data-driven, understandable and analysable, time-efficient, flexible, and ultimately... creative it was. If we do so, AI can be of enormous help in quantitatively supporting or rejecting decisions that need to be made using human intuition. The system is able to quantify these parameters, but interpretation remains the task of humans.

Cognitive support and visualization. We concluded the previous section by stating that interpreting measured and quantified values is a human task. We do not wish to refute this

idea, but it may be useful to examine how an AI-supported system can not only display data, but also significantly assist human thinking, situation recognition, and decision-making... with visual tools. This can reduce mental load and increase the speed of situational awareness and pattern recognition. AI does not think for humans, but – if used correctly – it can assist human thinking.

How can this be achieved? AI displays events in real time and, if necessary, can add verbal explanations (e.g., in 3D terrain, AR glasses, interactive maps, audio), thereby creating more transparent and understandable circumstances. Thus, the staff not only receive a set of data, but also see a picture of the situation in the present and the probable future. Therefore, visualization is not just an image, but also an act of communication.

Generative AI products are already being used to support the design and execution of wargames, making the wargaming itself more realistic on the one hand, and allowing additional material such as news, reports, images, and so on to be added immediately to the course of events on the other.

It must be able to suggest decision points, critical bottlenecks, and effects that may not be obvious at first glance, which the human mind or eye often fails to perceive. It is very important that the staff has the opportunity to ask “questions” and simulate assumptions with AI, for instance: “Show me what happens if the second battalion is delayed by 30 minutes.”

Visualization is not meant to make people think less, but rather to help them see more

clearly where and what is worth thinking about. The essence of AI support is not technological, but rather cognitive... like a “digital chief of staff” who keeps order among the information in the commander's mind. Visualization becomes truly cognitive support when people no longer look at the data but see and understand the meaning of the situation. The point, then, is that the wargame should “captivate” the participants as much as possible, allowing them to immerse themselves in it.

If the system is able to do this, the result will be a more transparent, intuitive analysis that will enable the commander to respond and make decisions more quickly and perhaps more accurately.

AI, as: wargamer, leader of the opposing force. An artificial intelligence-supported system must be capable of simulating realistic decision-making on the part of the enemy, rather than simply “playing” according to the doctrinal principles we have created. With the help of AI, it is possible – and necessary – to incorporate the personality traits of enemy commanders into behavioural models, but this should be done in a simplified, well-validated, and human-supervised manner. The goal is not to completely reproduce the human psyche, but to model decision-making preferences and reaction patterns for the purposes of wargames. Thus, instead of providing specific responses, AI is able to adapt flexibly to losses, deceptive manoeuvres, changes, and so on. In addition to all this, we must also mention one of the main characteristics of wargames, namely that playing a war-

game multiple times from the same starting position will always lead to different results, highlighting new opportunities and threats (UK Ministry of Defence, 2017). The system is able to continuously learn from these variations, essentially “training itself” like an artificial wargamer who becomes more prepared, more skilled, and gains more knowledge with each run. How does all this happen?

A well-calibrated virtual adversary commander can offer three major advantages to the wargame system. First, the staff will face a more realistic opponent, as the system does not repeat exactly the same pattern as a rule-based simulator would. The next thing is its ability to adapt much faster than humans, which allows it to react more quickly, thus posing a serious challenge to human commanders and staff. And finally, it should not be forgotten that systems operating in this way are constantly learning, so they become more accurate after each game because their experiences are continuously incorporated into their decision-making models.

Obviously, the question arises: if we try to model the decision-making of the opposing party – a human – as realistically as possible, who or what is best suited for this task? And this brings up countless other questions... A person who is an expert on the subject, or an AI wargamer? Can the knowledge of a Subject Matter Expert (SME) be replaced by an AI player? Are the decisions of the opposing party human at all, or do they rely heavily on AI suggestions? And so on... difficult questions, questions to which we can only guess at

the answers today, but at the dawn of this era, we do not yet have sufficient usable information about all of this. However, it already seems certain that the use of an AI wargamer can be useful when there are not enough participants with the right qualities, or when a substitute participant is needed to represent neutral actors (neutral forces, non-governmental organizations, or civilians).

One of the critical factors in the effectiveness of wargames is traditionally the ability and knowledge of the participants. However, in practice, SMEs on a given topic are often overburdened and not always available, which can be remedied, for instance, by the use of Large Language Models. In addition, the use of artificially generated data from wargames allows human players to learn about other perspectives, thereby discovering new possible solutions (Jensen, Atalan and Tadross, 2024).

Another crucial aspect – which generally applies to war as a whole – is that decision-makers must make decisions in an environment where not all information is available, or where there is a distinct lack of information, leading to what is known as the Fog of War. As Clausewitz expressed it: “...every war is rich in particular facts; while, at the same time, each is an unexplored sea, full of rocks, which the General may have a suspicion of, but which he has never seen with his eye, and round which, moreover, he must steer in the night.” (Clausewitz, Howard (ed.) and Paret (ed.), 1976, Book II, Chapter 2). Whether we use AI methods for scenario planning, decision support, as a wargamer, or

analysis, the fog of war can always arise. Managing this uncertainty is a key issue. One of the best solutions for this – which is used by most AI systems – is probability calculation. This makes uncertainty a little more tangible and quantifiable, which AI is able to handle, especially when large amounts of relevant data are available (University of Helsinki and MinnaLearn, 2018).

Nevertheless, the most worrying issue is not the functioning of AI, but humans themselves. This raises a very serious psychological and ethical question: what happens if leaders and staff begin to copy the machine and its decision-making? If this happens, it may even be that the work of staff becomes more efficient in theory, people may learn to make decisions faster, but they will consider less... they may even forget that warfare still takes place on the battlefield today and is also a very serious moral issue.

The authors do not wish to suggest that the use of AI is dangerous or should be avoided when simulating hostile activity. Not at all. Indeed, it is worthwhile and even necessary to exploit the enormous potential it offers... but this must be done sensibly and consciously. And in this case, awareness means control. In the future, AI systems should be operated with three restrictions:

- The first area is the “ethical filter” built into the AI system that “plays” the enemy. This filter must screen out suggestions that are humanly, legally, or politically unacceptable to one’s own armed forces or those of the opposing party. Creating and maintaining this is very difficult and not specifically a military

task. However, without it, the problems outlined above may arise during wargames, which will ultimately compromise the entire decision-making process. The question may also arise as to what a right is or just decision from a moral point of view. To what extent can such a subjective question be quantified? *Jus ad bellum* and *jus in bello* can be taught through machine learning... but how will they be applied in a wargame? In order to model reality with fidelity, we must also take into account the fact that the opposing party may not be bound by our moral standards. Therefore, it is important to emphasize the statement made a few sentences earlier, that the enemy must play by its own rules, even if they seem foreign to us.

- The following issue, as a continuation of the previous one, is the explainability of AI activities. This means that systems must be designed in such a way that AI is required to explain all of its decisions. Humans must be able to evaluate this explanation in every case and possibly reject it, even if it has already passed the first filter, the “ethical filter.” At this point, however, humans assign a task to AI on how to proceed from a certain decision point. This is a difficult but necessary question.

- Finally, barriers to the need for continuous human approval must be built into the system. This will likely slow down the process, but it is crucial to ensure that the actions of the AI-wargamer are truly “human.” However, it is critical to understand that the person approving decisions made by an AI tool must actually consider them and not purely rely on

them... just accepting them (Rivera et al. 2024).

And lastly, a highly philosophical question that is still a long way off. An AI commander is not a “real” opponent until it has intent. Intent, however, requires many human qualities, such as a value system. A human commander makes decisions because they want to achieve something. An AI probably makes decisions because it wants to maximize something. It would truly become a “commander” if it were able to understand why it wants to win, but that may go beyond the question of artificial intelligence and into the realm of artificial will... which is probably already at the level of artificial superintelligence?

There is another concern regarding the implications of using AI-based systems as players in wargames or to support decision-making. Hunter and Bowen argue that current narrow AI systems based on machine learning are fundamentally unsuitable for the command decision-making required in warfare, as their operation is based on inductive logic – that is, the recognition of past patterns – whereas the essence of decision-making from the tactical and the strategic levels is abductive reasoning, i.e., decision-making characterized by uncertainty and the fog of war (Hunter and Bowen, 2024).

To this add that models can study, for instance, Clausewitz’s work or other military thinkers’ “timeless” works, but will they be able to display the phenomenon of “military genius” described by him during a wargame? We do not know. Can models replace real

military experience gained “firsthand” with probability calculations and conclusions based on large data sets? We cannot know this for sure either...

*Analysis, evaluation of results, or learning
and lessons identified and learned*

At its current level of development, an AI system is able to record the results of hundreds or even thousands of previous wargames in a database, and with the help of machine learning, it refines its internal worldview with each run and each replay, increasingly better estimating what the outcome of certain situations might be. In addition, it recognizes which activities, patterns, and doctrinal procedures lead to success or failure and reduces the randomness of its decisions. We could also say that it becomes increasingly “confident” because unfamiliar situations will increasingly resemble something it has seen before, and it will make similar decisions based on identical inputs. It can then use this growing “knowledge” in subsequent wargames.

Therefore, unlike human staff, the more we “play” the system, the smarter and more sophisticated it becomes... instead of becoming tired. This sounds very nice, but is it really true?

Indeed, there are several issues. The first is mathematical in its aspect. The development presented above is statistical in nature... current AI presumably does not understand what it is learning, but rather its knowledge increasingly accurately matches past patterns, so the “intelligence” of AI is not a deepening

of thought, but rather a refinement of statistics. However, it logically implies that machines will become better at data, while humans will become better at interpreting things.

The second area to consider: when talking about AI systems “learning” is misleading. Human learning often develops through mistakes, irregularities, and surprises. AI cannot “learn” from mistakes in the sense of understanding why a mistake occurred... it only statistically adjusts the metrics so that it does not happen again. Machines learn from repetition, humans learn from interpreting contradictions. The system develops in terms of accuracy, while humans – hopefully – develop in terms of meaning.

The third area of concern is “mental state.” AI can run continuously, never gets bored, never gets stressed, never changes its strategy... based on mood. Although this requires significant resources. On the one hand, this is an advantage, but on the other hand, it is also a cognitive disadvantage, because human fatigue, doubt, and mood swings are not mistakes, but learning fluctuations: these are the moments when the decision-maker breaks their routine. The machine, on the other hand, never questions its own logic, it only refines it. The commander, however, sometimes consciously does not believe in statistics, and this is precisely what creates the unexpected.

Overall, it can be concluded that the more we use the system supported by AI, the more it develops. It develops because it creates more accurate patterns and makes more stable, reliable predictions. However, it is

almost certain that AI systems at their current level will not become more intelligent in the human sense, because they do not know the “why,” only the “how”; they do not learn to think, they only become better at calculating. And finally, they will not become wiser, because wisdom is the ability to recognize contradictions, which machines are likely to see as errors.

An AI method can therefore learn about combat and armed conflict, but it does not actually learn how to wage war. Humans learn to wage war but often forget what they previously knew about war. The more we “play” the system, the more accurate its probability estimates become, the more consistent its behaviour becomes, and thus the less novelty it will bring, because it has already seen everything that its model allows. However, if we try to teach the machine not only the results but also the reasons behind human decisions, then the system not only “learns” but also “reflects,” and this comes close to true artificial intelligence. The two – human and machine – can become truly powerful together, because then the AI remembers and the human interprets.

Although professional military wargaming is the focus of this paper, it is worth mentioning that this topic does not only cover the issue of warfare, as for instance deescalation, peace operations, and human security are equally important from a military perspective. Therefore, the military application of wargames is not strictly limited to combat. For instance, a wargame focused on peace-making requires a completely different logic, set of criteria, and

decision-making process than, say, “testing out” course of actions. In the former case, the humanitarian focus, social dynamics, and a broad interpretation of security come to the fore, while in the latter case, the emphasis will be on the application of doctrinal principles and the practice – and, if necessary, modification – of combat-level procedures. Based on this, the AI tool used will also need different data and inputs. It follows that AI-based solutions cannot be applied uniformly in the same way in every military exercise but must be adapted to the context and the desired outcome.

The benefits of using AI

When we look at modern armed conflicts, we see that the speed of events has long since exceeded the pace that the staff can keep up with manually, using their creative and analytical minds. The speed of information flow is usually measured in fractions of a second, not minutes or hours. A human staff is simply not capable of processing and interpreting this amount of data. This is a situation we must accept and try to deal with. Restructuring staff and increasing their numbers is only a solution to a certain extent. The use of modern communication tools and computer technology is only an aid; it does not fundamentally solve the problem. AI is therefore probably not a luxury, but a necessary neurological aid, or, so to speak, an improvement to the central nervous system of the armed forces. If we do not take this into account and perhaps decide not to use AI, it is not a morally superior

approach, just a slower one... In the wars of the future, it is certain that the speed of thought will be the fire-power itself.

The current culture of wargames has been around for decades, even centuries, because the human mind could still “see through” the system: a few units, a few manoeuvres, predictable intensity, known enemies. The wargames of the 18th and 19th centuries took place on the board, those of the 20th century on the map, and those of the 21st century increasingly in the data space... This is not a question of romanticism; science always develops with the help of the tools available. Today, some types of wargames are – for the most part – too complex to be held together by human imagination alone. AI is the only tool capable of handling statistical uncertainty and logical coherence at the same time at this scale. Although we have to note that this cannot be applied to all types of wargames, for instance, the rules and mechanisms of matrix wargames are less rigid and rather enhance on creativity and brainstorming.

It should be noted, however, that there is another side to this issue. The push towards digitization may overlook the human aspect. Analog wargames will continue to be used. These are important because, among other things, they have advantages that no software seems to be able to replace. It is more enjoyable for players to be able to touch the pieces and markers, as this allows them to participate actively in the course of the war-game, immerse themselves in it more deeply, and place greater emphasis on personal interaction.

And although the battlefield has always been multidimensional – only the number of dimensions varied in different era – today, an unmanageable amount of data is generated in every “dimension.” The kind of intuition and insight that comes from experience, which we have interpreted as military genius in almost every era, is no longer enough. It is necessary, but not sufficient for success...

In earlier periods of warfare, individual commanders could exert a decisive influence, even though war was always complex. Today, however, due to complexity and vast data flows, quantitative processing itself has become a strategic capability. A wargame can focus on any time period in the past or distant future. For instance, it can explore a past conflict for educational purposes, or it can focus on the near future to support decision-making in relation to a current conflict. A well-planned and well-designed wargame – if it wants to and is able to predict the future – has to run through millions of combinations. Humans cannot do this because their attention, time, and energy are limited. AI, on the other hand, is tireless, or at least much less prone to fatigue than humans... we do not know this for sure yet. It is not needed because it has a better understanding of military science, but because it is able to perform repetitive thinking tasks that humans cannot do at the speed required today. This frees human decision-makers from the drudgery of analysis and thus truly supports us in our decision-making.

Indeed, this will likely result in a completely new situation, as there is no reason to assume

that if we use AI, our enemies will not do the same. Therefore, in wargames, it is not enough to model human opponents – we also need simulations of AI against AI. So, in the 21st century, human commanders are not only fighting against the human decisions of the enemy, but also against the evolution of machine strategies.

Let us now return to humans for a moment. Humans tend to forget their mistakes. This is even more true in military institutions... failures are rarely part of the curriculum, and in many cases are almost a political burden, something to be forgotten. AI, on the other hand, is able to learn from every mistake, since every wrong decision or incorrect reaction is actually data. Wargames can be built on mistakes, which are not failures but feedback – unwanted outcomes, deviations in the process. This is not very appealing, but it is certainly useful in the long run. Military science cannot renew itself if the past must be hidden and ashamed of. AI is not ashamed; rather, “it” archives.

At this point, we have arrived at an important aspect of wargaming. During wargaming, the emphasis is not on winning at all costs and ensuring that the opposing side loses, but rather on how and why we did not win, which is often a valuable result in itself (Harangi-Tóth, 2020). What were the factors, decision points, and shortcomings that led to “not winning”? And it does all this in a safe-to-fail environment where we can learn even from the outcome of a losing without any real losses.

And now consider another human concern... Military history is replete with tragedies caused by the overconfidence of commanders. The human mind loves to believe that it can control chaos. The introduction of AI will put an end to this illusion, and humans will be forced to admit that they do not fully understand the decisions made by complex systems. Understanding this is essential for understanding and controlling the battlefields of the future. However, once we understand this, AI will not take away human power but rather realize its limitations and make it more sustainable. In a world where war can paralyze global infrastructure in minutes, excessive human overconfidence is a greater danger than machine coldness.

The integration of artificial intelligence should not be about machines making decisions and humans accepting them, but rather about combining the two forms of thinking mentioned earlier, analytical and intuitive, into a single system. And this is where wargames become really important, as this is the arena where this cooperation can be safely experimented with and then practiced. While machines learn from human decisions, humans learn from machine logic. This is not subjugation on either side, but a kind of alliance, let’s call it a cognitive alliance, perhaps the first truly new type of human-technology relationship in strategic thinking. This may be an accurate statement, as more and more people today – including the authors themselves – are insisting that the domains should be supplemented with a new one: the cognitive domain.

We believe that if we do not integrate AI, we will not preserve humanity or human dominance, but rather human limitations. And history has always shown that those who cling excessively to their limitations will sooner or later perish among them.

Limitations and risks

AI – as it currently stands – appears to be able to help us cope with the incredible dynamics of the battle and the enormous amount of data. But before we start thinking that this will solve all our problems at a stroke, we must also mention the expected downsides.

The above described factors provide significant assistance in the planning, design, execution and analysis of wargames played within the framework of military decision-making. Nevertheless, we should not ignore certain concerns. These concerns stem from the most fundamental differences between humans and AI. An AI “wargamer” is not human, so it may interpret various basic concepts, such as “victory,” differently. At the current level of development of these systems, it is safe to say that their decisions, or the decisions they recommend to humans, do not yet have a moral dimension; that is, they do not understand the concept of “too high a price,” for instance. It is possible that if the probability of success increases, AI would be able – at least today – to make (or recommend) decisions that are unacceptable to humans. It will probably not fear loss, it will have no emotional barriers, so it may be inclined to

follow the principle of “maximum risk, maximum profit,” which human commanders – without emotion – are incapable of. However, this means that simulating its activity cannot be realistic either. For an AI commander, war, combat, killing, and loss are not yet a tragedy, so AI will probably always be logical, but not always reasonable... and the difference between the two would be measured in human lives on the real battlefield.

AI-based tools are able to process data and information and performing analysis and evaluation at unprecedented speeds and volumes, but they can also potentially slow down decision-making, as they generate even more data and information for decision-makers to consider, raising further questions (Carnegie Endowment for International Peace, 2024).

When using AI tools in wargames – as a wargamer or decision support – it is also important to consider how this might affect the very human group dynamics. During wargames, participants find themselves in decision-making situations. Imagine, for example, that in such a decision-making situation, participants have to come up with a solution, and an AI “player” also comes up with a proposal. To what extent will participants be inclined towards the AI’s point of view and proposal as a result of excessive trust and reliance on technology? And if one of the experts on the subject, even if they have a different opinion, questions the solution proposed by the AI, will they do so? It is more than likely that when a group has to make a

decision in a short period of time, there will be a strong urge for participants to follow the recommendations of the AI “player” (see: Carnegie Endowment for International Peace, 2024).

Perhaps an even bigger concern is excessive reliance on AI, where the staff members uncritically accept AI outcomes and recommendations, thereby eliminating creative and critical thinking. If the staff rely too heavily on the “wisdom” of AI systems, cognitive courage, the kind of intuitive boldness that is often at the heart of critical decision-making, may indeed decline.

In wargames, AI is not primarily a simulation tool, but rather an aid to human decision-making, which is also an influencing factor. For this reason, however, any errors, distortions, or misinterpretations can have serious consequences. If AI recommends a course of action in a way that cannot be understood, explained, or analysed, the command staff will likely treat it with caution. In the case of AI systems, interpreting natural human language can also pose problems. For instance, due to multiple meanings, AI systems may encounter difficulties in correctly interpreting the instructions they receive. However, research is already underway to address this issue (Schadd, 2022).

The accuracy of the data entered can also be a serious source of concern. If the initial data is incorrect or does not fully reflect reality, the system's simulation, answers to questions, and recommendations will also be distorted.

Finally, there is a lack of creativity. Current systems think in terms of learned patterns and

do not “think outside the box.” This means that they will probably continue to lack intuition for a long time to come. However, if AI learns from its own past, sooner or later its worldview will become closed, so the algorithm will learn from the same patterns over and over again and preserve its thinking logic, which poses two dangers.

AI starts from existing patterns and data and is likely to consider those successful that previous human staff considered successful in their staff work. In doing so, the system effectively elevates the doctrines of the past to the norm and interprets true novelty as a statistical “anomaly.” The other danger is a bit more philosophical, but it’s worth thinking about. If an AI-supported system works from the past and isn’t prepared for the future, it’ll probably know everything about what’s already happened and less about what never will... At the same time, it can recognize trends on a statistical basis and making predictions based on probability, but it does so based on existing data, so it will not be able to incorporate data it does not yet have into its probability calculations.

Machine learning involves processing extremely large amounts of data, which ultimately – at least for now – originates from the human world, meaning that the data sets inevitably contain distortions that will also be evident in AI models, producing discriminatory features (Leavy, O’Sullivan and Siapera, 2020). This carries numerous risks; just think of a wargame with a geopolitical theme. What negative impact could it have on its fidelity and results if, for instance, the AI

tool adopts a Western-centric (or any other) worldview because that was predominant in the data set?

Another issue related to machine learning is that it can be an extremely costly process, especially if we want to create a complex system, while in wargaming it needs to be flexible and modifiable so that it can be used in other scenarios or for other purposes. For this reason, it may be preferable to use „smaller” modular AI systems in wargames (Ryseff and Bond, 2022).

A 2024 study examined the military and diplomatic decision-making behaviour of five large language models – extending to the use of nuclear weapons – within the framework of wargames, which yielded extremely thought-provoking results. All of the models examined were more prone to escalation – to varying degrees – than human participants, and this often manifested itself in the form of sudden escalation that was difficult to predict. In addition, they tended to provoke an arms race... both conventional and nuclear. With all these decisions, the models typically justified deterrence and first-strike tactics. One possible explanation for this tendency toward escalation, according to the authors, is that publications related to escalation predominate in the international literature, so the models, learning from the written material, adapted this focus (Rivera et al. 2024). The escalatory behaviour of the models suggests that decision-making based solely on statistics and probability calculations is insufficient, as non-material aspects such as human life must also be taken into account.

This tendency toward escalation also poses a risk when AI is used to support decision-making, as it can influence the decision-maker’s views in an unintentional and subtle way, and raises further questions about the logic behind decisions made by such models, which is not entirely transparent (Barzashka, 2023). All of this once again highlights the need for an “ethical filter,” explainability, and human approval with critical thinking in mind. As with any new innovation or tool, we must ask ourselves to what extent we want to rely on it. There is a noticeable trend towards the digitization of wargames and, in parallel, AI-supported wargaming, thus marginalizing manual wargaming, even though, as already mentioned, it also has its advantages. If we rely too heavily on AI-supported processes and decision-making, how will this affect resilience if, for instance, an unexpected event causes the AI-based tool to become temporarily or permanently inoperable...

When weighing up the limitations and risks, we may be inclined to conclude that this involves too much risk, but we should also consider what might happen if we neglect such an innovative development. When exploring the integration of AI tools into wargames, it is worth bearing in mind that developments in this area are not limited to NATO member states. The People’s Republic of China has made significant capability developments in both AI and wargames in recent years and continues to do so today. For instance, AI systems are being used to develop wargames with complex scenarios (Black and Darken, 2023). Consequently, both at NATO

and national level, developments and research in this area are essential in order to the alliance to maintain its competitive edge over potential adversaries. Failure to participate in such innovative processes could potentially result in falling behind.

There is lively professional debate around the world on how to mitigate or completely eliminate these limitations and risks. Some experts see hierarchical reinforcement learning as the solution. This is an extension of traditional reinforcement learning, which attempts to address limitations by introducing a hierarchical structure into the learning and decision-making process. Its central element is that complex, longterm tasks are broken down into smaller, more manageable subtasks in a hierarchical system. As part of this, the AI system can apply multiple models depending on the situation, thereby optimizing decision-making and action strategies (Black and Darken, 2023).

Further research will definitely be needed in terms of limitations and risks so that AI systems can be used in the most optimal and innovative way possible in the field of wargaming.

Summary

– new dimension and looking ahead

AI is beginning to have an impact on warfare today, and it is almost certain that it will generate very significant changes in the near future, along with other EDTs. For this reason, the new systems are sure to have a significant impact on military operations as well.

In the foreseeable future, armed conflicts will still mainly involve human actors, so their main characteristics will remain virtually unchanged, despite ongoing minor changes. In the short term, machines will not take over or “want” to take over the role of humans in plan-ning, organizing, decision-making, or execution, so for the time being, we must not forget Clausewitz’s basic principles. However, there is a real danger that, despite the fact that machines do not yet have a will of their own, we ourselves will unwittingly hand over some of our roles and responsibilities to them... because it is more convenient, faster, and more precise. This is not necessarily a bad or harmful process, but – as in so many areas of life – a healthy balance must be found here too. If we fail to do so, we may end up living with a perfectly optimized system that knows everything about warfare except one thing: what people do and why, why people wage war...

The next issue concerns the very existence of AI. The roots of research in this field actually go back to Alan Turing, and since then it has experienced countless ups and downs. There was a period that we now refer to as the AI winter, but today we are witnessing and participating in an unprecedented boom. It has become fashionable not only to work on, re-search, and develop AI, but also to talk about it. We ourselves are doing just that now... However, it is worth considering whether this research and development will yield any serious results. Is artificial super-intelligence, or even AGI even achievable in the near future? Could there be another AI

winter, or, on the contrary, can we expect further major breakthroughs? One thing is certain – as already mentioned – the capabilities of existing AI systems already have a significant impact on warfare, and AI research will be decisive for future warfare.

Decision-making, and its speed and accuracy, is always a crucial issue in warfare. In developing this, innovative approaches such as the use of AI systems in wargames, where decision-making is the central element, are unavoidable. As has been demonstrated, AI-based solutions can effectively and in many ways assist in the planning, designing, execution, and analysis of wargames, while also accelerating these processes.

A key question is how this can be implemented most effectively and what the first steps of initial implementation should be. The authors of this study believe that it is worthwhile to first integrate AI-based solutions into smaller-scale wargames, the results of which can then be incorporated into more complex wargames. and at the domestic level, it is definitely recommended to monitor efforts in this direction within the alliance and to learn from them, adapting them to our culture if necessary.

As emphasized several times in the paper, the use of AI-based solutions in wargames should only be introduced with the restrictions discussed, i.e., taking ethical considerations into account. Particular attention must be paid to ensuring that the AI system is always able to explain its decision proposals, that humans can follow the steps leading to the proposal, and that human approval is an

integral part of the process. These are unlikely to be easy tasks, but they are essential for effective and efficient application in order to reduce or even eliminate the risks described above.

The authors believe that, alongside digitization and the use of AI systems, we must not forget about analogue wargames, their advantages, and the role of humans. From what has been described so far, it is clear that we are still a long way from the optimal situation, so we recommend further research into wargames and AI systems, with a particular focus on mapping the development of AI systems and their impact on warfare, for instance.

In conclusion, let us assume that AI tools will be successfully and effectively integrated into the wargaming procedures of our own and our allies in the near future. With this assumption, we can also be sure that potential adversaries are researching and developing this capability, so it is likely that if we succeed, they will too. Innovative developments of this kind certainly take time – in terms of procedures, technology, and personnel – so if their development does not begin in time, it cannot be made up overnight. Falling behind is not an option, either domestically or within the framework of the alliance in general.

The integration of AI into wargames may carry potential risks, but its absence will clearly lead to disadvantages. The future of warfare is not a question of human or machine, but whether we are able to elevate thinking as a biological function to a technological capability and whether we will be able to apply

innovative solutions effectively... The successful strategist of the future will no longer be the one who is good at math, but the one who is able to think together with an intelligent system while preserving the one thing that machines – it seems – will not learn for a very long time:

the shame of responsibility.

References

- Barzashka, I. (2023). *Wargames and AI: A dangerous mix that needs ethical oversight*. Bulletin of the Atomic Scientists, 4 Dec 2023. URL: <https://thebulletin.org/2023/12/wargames-and-ai-a-dangerous-mix-that-needs-ethical-oversight/> [Accessed 18 November 2025]
- Black, S. and Darken, C. (2023). Scaling Artificial Intelligence for Digital Wargaming in Support of Decision-Making. In: *STO-MP-MSG-207: Proceedings of the 207th Meeting / Symposium*, NATO Science and Technology Organization. URL: <https://publications.sto.nato.int/publications/STO%20Meeting%20Proceedings/STO-MP-MSG-207/MP-MSG-207-23.pdf> [Accessed 16 November 2025].
- Caffrey, M. (2019). *On Wargaming. How Wargames Have Shaped History and How They May Shape the Future*. Newport: U.S. Naval War College. URL: <https://digital-commons.usnwc.edu/newport-papers/43/> [Accessed 16 November 2025].
- Carnegie Endowment for International Peace (2024). *How AI Might Affect Decisionmaking in a National Security Crisis*. URL: <https://carnegieendowment.org/research/2024/06/artificial-intelligence-national-security-crisis?lang=en> [Accessed 14 November 2025].
- Clausewitz, C. (1832). *On War*. Trans. Howard, M. and Paret, P. Princeton University Press, 1976.
- Doktrinzentrum der Bundeswehr (2024). *Wargaming Handbuch der Bundeswehr*. Bundeswehr Verlag. URL: <https://www.bundeswehr.de/resource/blob/5872302/f33665313cbbf07c098ed7e3769530ae/handbuch-wargaming-bundeswehr-data.pdf> [Accessed 13 November 2025].
- Harangi-Tóth, Z. (2019). *Mi a történelmi hadijáték és miért van helye a katonai felsőoktatásban?* In: *Hadtudomány*, XXIX (4), 119–128. Doi: <https://doi.org/10.17047/HADTUD.2019.29.4.119> [Accessed 12 November 2025].
- Harangi-Tóth, Z. (2020). Nemzetközi hadijátéktrendek. *Honvédségi Szemle*, 148 (2), 113–124. Doi: <https://doi.org/10.35926/HSZ.2020.2.11> [Accessed 13 November 2025].
- Hunter, C., and Bowen, B. E. (2024). We'll never have a model of an AI major-general: Artificial Intelligence, command decisions, and kitsch visions of war. *Journal of Strategic Studies*, 47(1), 116–146. Doi:

- <https://doi.org/10.1080/01402390.2023.2241648> [Accessed 23 November 2025].
- Jensen, B., Atalan, Y. & Tadross, D. (2024). *It Is Time to Democratize Wargaming Using Generative AI*. Center for Strategic and International Studies, 22 February 2024. URL: <https://www.csis.org/analysis/it-time-democratize-wargaming-using-generative-ai> [Accessed 30 November 2025].
- Leavy, S., O'Sullivan, B. and Siapera, E. (2020). *Data, Power and Bias in Artificial Intelligence*. Doi: <https://doi.org/10.48550/arXiv.2008.07341> [Accessed 21 November 2025].
- NATO (2022). *NATO 2022 Strategic Concept*. URL: <https://www.nato.int/strategic-concept/> [Accessed 07 December 2025].
- NATO Allied Command Transformation (2023). *NATO Wargaming Handbook*. Virginia. URL: <https://paxsims.files.wordpress.com/2023/09/nato-wargaming-handbook-202309.pdf> [Accessed 07 November 2025].
- Petrovánszki, J. (2023): *Hadijátékok aktualitása a jelenlegi biztonságpolitikai környezetben, valamint szerepük a stratégiai elemzésben és gondolkodás fejlesztésében*. Bachelor's thesis. University of Public Service
- Pielström, S. and Wintjes, J. (2019). Preußisches Kriegsspiel: Ein Projekt an der Julius-Maximilians-Universität Würzburg. *Militärgeschichtliche Zeitschrift*, 78(1) 86-98. Doi: <https://doi.org/10.1515/mgzs-2019-0004> [Accessed 13 November 2025].
- Rivera, J.-P., Mukobi, G., Reuel, A., Lamparth, M., Smith, C., and Schneider, J. (2024). *Escalation Risks from Language Models in Military and Diplomatic Decision-Making*. In: 2024 ACM Conference on Fairness, Accountability and Transparency (FAccT 2024), pp. 836–898. Doi: <https://doi.org/10.1145/3630106.3658942> [Accessed 28 November 2025].
- Ryseff J, Bond M. Small is beautiful. *The Journal of Defense Modeling and Simulation: Applications, Methodology, Technology*. 2022;22(1):19-23. Doi: <https://doi.org/10.1177/15485129221096478> [Accessed 22 November 2025].
- Sabin, P. (2014). *Simulating War*. Bloomsbury Academic.
- Schadd M, Sternheim AM, Blankendaal R, van der Kaaij M, Visker O. How a machine can understand the command intent. *The Journal of Defense Modeling and Simulation: Applications, Methodology, Technology*. 2022;22(1):41-58. Doi: <https://doi.org/10.1177/15485129221115736> [Accessed 18 November 2025].
- Somkuti, B. (2021). *Ólomkatonák, avagy a hadijátékok története*. Index.hu, 14 April 2021. URL: <https://index.hu/techtud/tortenelem/2021/04/14/haborus-jatekok-a-wargame-tortenete-jatekkatona-olomkatona-tortenelem/> [Accessed 01 November 2025].
- Szinonimák.hu (online synonym dictionary). „hadijáték”. URL: <https://szinonimak.hu/hadij%C3%A1t%C3%A9k/>

- [C3%A9k-szinonima](#) [Accessed 01 November 2025].
- UK Ministry of Defence (2017). *Wargaming Handbook*. Ministry of Defence (DCDC), 4 August 2017. URL: https://assets.publishing.service.gov.uk/media/5a82e90d40f0b6230269d575/doctrine_uk_wargaming_handbook.pdf
- University of Helsinki & MinnaLearn (2018). *Elements of AI*. Online course (MOOC). URL: <https://www.elementsofai.com/>
- U.S. Military Academy Library (USMA Library). *Generative Artificial Intelligence (GenAI) – Capabilities and Resources*. URL: <https://library.westpoint.edu/GenAI> [Accessed 12 November 2025]

**A BIZTONSÁG MEGÍTÉLÉSÉNEK VIZSGÁLATA MOZGÁSSÉRÜLT VEZETŐK
ESETÉBEN – A MESTERSÉGES INTELLIGENCIA SZEREPÉNEK
KITERJESZTÉSÉVEL**

Author(s) / Szerző(k):

Adrienn Oravecz (PhD)

Semmelweis University

(Magyarország)

E-mail:

oravecz.adrienn@semmelweis.hu

Cite: Oravecz, Adrienn (2025): A biztonság megítélésének vizsgálata mozgássérült vezetők esetében – a mesterséges intelligencia szerepének kiterjesztésével. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 35-44.

Idézés:

Doi: <https://www.doi.org/10.35406/MI.2025.2.35>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0008

Reviewers: Anonymous reviewers / Anonim lektorok:

Lektorok:

1. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
2. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

A tanulmány célja a közlekedésbiztonság megítélésének vizsgálata mozgássérült vezetők körében. Jelen tanulmány keretein belül, szeretnénk bemutatni azt is, hogy jelenleg is már hogyan járul hozzá a mesterséges intelligencia (MI) közlekedésbiztonsághoz az autóvezetés során. Valamint, hogy milyen potenciál rejlik benne még a jövőben. Az empirikus kutatás során 61 fő mozgáskorlátozott személy vett részt online kérdőív segítségével. A kereszttábla-

elemzések azt mutatták, hogy a rendszeresen vezető mozgássérültek 94%-a biztonságban érzi magát vezetés közben, ami a vezetési rutin és a szubjektív biztonságérzet közötti pozitív összefüggést jelzi. A mesterséges intelligencia által nyújtott vezetéstámogató rendszerek és az önvezető járművek új lehetőségeket kínálnak a közlekedésben való részvétel kiszélesítésére, ugyanakkor etikai és infrastrukturális kihívásokat is felvetnek. A tanulmány rávilágít arra, hogy a biztonságérzet nem csupán technikai, hanem pszichológiai és társadalmi dimenziókat is magában foglal (Zeng és mtsai., 2024; Patil és mtsai., 2024; Lv és mtsai., 2023).

Kulcsszavak: közlekedésbiztonság, mesterséges intelligencia, mozgássérült vezetők, önvezető járművek, defenzív vezetés, Munsch

Diszciplínák: neveléstudomány, pszichológia

Abstract

EXAMINING THE PERCEPTION OF SAFETY AMONG DRIVERS WITH PHYSICAL DISABILITIES – WITH AN EXTENDED FOCUS ON THE ROLE OF ARTIFICIAL INTELLIGENCE

The aim of this study is to examine the perception of road safety among drivers with physical disabilities. Within the framework of this paper, we also seek to present how artificial intelligence (AI) currently contributes to traffic safety in driving, as well as the potential it holds for the future. The empirical research involved 61 participants with mobility impairments who completed an online questionnaire. Cross-tabulation analyses showed that 94% of drivers with disabilities who drive regularly feel safe while driving, indicating a positive relationship between driving experience and perceived safety. AI-based driver assistance systems and autonomous vehicles offer new opportunities for expanding participation in transportation; however, they also pose ethical and infrastructural challenges. The study highlights that the sense of safety encompasses not only technical but also psychological and social dimensions (Zeng et al. 2024; Patil et al. 2024; Lv et al., 2023)

Keywords: road safety, artificial intelligence, drivers with disabilities, autonomous vehicles, defensive driving, Munsch

Disciplines: pedagogy, psychology

A közlekedésbiztonság fogalma az utóbbi években jelentős változáson ment keresztül, különösen a mesterséges intelligencia (MI) térhódításának hatására. A hagyományos megközelítések a balesetek megelőzését és a fizikai

biztonsági intézkedéseket helyezték előtérbe, míg napjainkban a humán tényezők, a vezetési kultúra és a technológiai támogatás integrált szemlélete vált meghatározóvá. A mozgássérült vezetők esetében a biztonság kérdése

különösen fontos, mivel a fizikai akadályok mellett gyakran pszichológiai és infrastrukturális korlátok is nehezítik a közlekedésüket. Jelen kutatás célja, hogy feltárja, hogyan ítélik meg a biztonságot a mozgássérült vezetők (Khan és Das, 2024; Booker és mtsai, 2024; ITF, 2022, 2023).

Közlekedésbiztonság

A közlekedésbiztonság meghatározására jelenleg nincsen egy általánosan elfogadott definíció a 2020-2025-ben között megjelent szakcikkekben. Meghatározása erősen kontextusfüggő, de egy tendencia látszik kirajzolódni, amely törekszik, arra, hogy több tényezőt és nézőpontot vegyen figyelembe. Így adva lehetőséget egy komplexebb értelmezés kialakítására és használatára (Khan és Das, 2024; Booker és mtsai, 2024; ITF, 2022, 2023).

A közlekedésbiztonság fogalmát a WHO 2023-as írása a balesetekből eredő sérülések és halálozások csökkentésének komplex rendszerként határozza meg, amely a prevenciótól az infrastruktúra tervezésén át az utólagos beavatkozásokig terjed. A 2020-as évektől kezdődően a „Safe System” megközelítés vált uralkodóvá, amely nem a hibás egyént, hanem a teljes közlekedési rendszert tekinti a biztonság zálogának (Khan és Das 2024; Booker és mtsai 2024; ITF, 2022, 2023).

Műszaki és közlekedéstervezési felfogásban a biztonság elsősorban az infrastruktúrát és járműszintű beavatkozásokat foglalja magában.

Magatartás/pszichológia orientált definíciók: Munsch kutatásaira alapozva a közlekedésbiztonság a humán tényezőket (figyelem, alkohol, kockázatvállalás) és a kultúrát (safety culture) hangsúlyozzák. A rendszer minden eleme – utak, járművek, sebesség, emberi magatartás – összefüggésben áll, így a felelősség megosztott. Gerhard Munch lényegében a „defenzív vezetés atyja” (Közlekedésbiztonság 2021.08.16.). Munsch nevét, a forgalom- és közlekedépszichológia korai németországi kezdeményezőjeként említik. Már az 1960-as évektől kezdve foglalkozott kezdő sofőrökkel és a vezetés tanításának észlelési/megfigyelési aspektusaival. A „forgalmi látás” és veszélyelőrejelzés kutatása fűződik nevéhez. Munkássága a vezetés gyakorlati oktatásában máig jelentős. A vezetőoktatás, „defenzív taktika” és forgalmi ismeretek tantervi megközelítései — munkái hatottak a forgalmi ismeretek oktatásának gyakorlatára. Például, Svájcban használt anyagok eredete részben hozzá köthető. Munsch az emberi tényezőt kezdte el hangsúlyozni a közlekedésbiztonságban – tehát nem csak a jármű- vagy úttechnikai oldalt, hanem a pszichológiai és viselkedési aspektusokat is.

A közlekedésbiztonság vizsgálata globális szinten is egyre inkább multidiszciplináris jellegű. C. Pieres és munkatásainak 2025 írása kiemeli, hogy a közlekedésbiztonsági teljesítmény összehasonlítása csak egységes indikátorrendszeren keresztül lehetséges. A baleseti statisztikák, vezetői viselkedési mutatók, infrastruktúra, szabályozási környezet és a közlekedési kultúra együttes értékelése olyan komplex képet ad, amely segíti a dön-

téshozókat a biztonsági stratégiai fejlesztésekben. A kutatás szerint a magas közlekedésbiztonsági szintű országokban erős a megelőzésre épülő szemlélet, a szabálykövetés kultúrája, illetve az infrastruktúra fejlettsége.

Du és munkatársainak 2023-as kutatása rávilágít arra, hogy a vezetői viselkedés elemzése napjainkban kiemelt szerepet kap a közlekedésbiztonság előrejelzésében. A hagyományos baleseti adatok mellett egyre nagyobb hangsúlyt kapnak az olyan viselkedési mutatók, mint a gyorsítási mintázatok, hirtelen fékezések, sávelhagyások vagy a vezetői reakcióidők. A telematikai eszközökből, szenzorokból és AI-alapú rendszerekből származó valós idejű adatok lehetőséget adnak arra, hogy korábban nem látható kockázati mintázatok váljanak azonosíthatóvá. A mesterséges intelligencia és gépi tanulási algoritmusok alkalmazása a viselkedésalapú kockázatbecslésben nagy pontosságot biztosít, és hozzájárul ahhoz, hogy a közlekedésbiztonsági beavatkozások célzottabban valósuljanak meg.

E két kutatás együttesen rámutat arra, hogy a közlekedésbiztonság értékelése a jövőben egyre inkább a humán, technológiai és társadalmi tényezők integrált vizsgálatán alapul, amely különösen fontos lehet a mozgássérült vezetők közlekedési részvételének megértésében és támogatásában.

Az MI-alapú vezetéstámogató rendszerek és az önvezető járművek fejlődése új távlatokat nyit a közlekedésbiztonságban, különösen a fogyatékkal élők számára (Zeng és mtsai 2024).

A mesterséges intelligencia közlekedésbiztonságban betöltött szerepe a mozgássérült vezetők szempontjából

Az „MI alapú” (mesterséges intelligencia) rendszerek még nem feltétlenül célzottan csak mozgássérült vezetőknek készülnek, de jelenleg is már több olyan technológia létezik vagy fejlesztés alatt áll, amelyek jelentősen növelhetik a biztonságot és az akadálymentesebb vezetési élményt számukra. A teljesség igénye nélkül ismertetünk most néhányat ezek közül jelen írás keretei között (Andrade és mtsai 2024; Liu és mtsai, 2023).

1. Fejlett vezetősegítő rendszerek (ADAS – Advanced Driver Assistance Systems). Az ADAS rendszerek – pl. sávelhagyás-figyelés (Lane Departure Warning), sáveltartás (Lane Keeping), adaptív sebességtartó automatika (ACC), vészfékezés – már széles körben elérhetők. Ezek a technológiák „intelligensen” dolgoznak szenzorok, kamerák és algoritmusok segítségével, hogy csökkentsék a vezető terhelését.

2. Személyre szabott vezérlés / adaptív segédrendszerek sérült vezetők számára. Elkezdtek olyan kutatások, amelyek elektromiográfiát (sEMG) használnak: az izmok elektromos aktivitását mérve a rendszer tud adaptív beavatkozást biztosítani a kormányzásban, például, ha a vezető csak egy karját tudja használni (Zeng és Liu 2024)

3. Személyfigyelés (Driver Monitoring Systems, DMS). A CogniSafe egy mesterséges intelligenciával támogatott rendszer, amely deep-learning és számítógépes látás (computer vision) segítségével figyeli a vezető

állapotát. Például, figyelem, fáradtság, elkalandozás vezetés közben (Schmidt és Weber 2023).

4. Hangvezérlés és másodlagos funkciók. A Mobility in Motion 2022 kínál olyan vezetéssel segítő eszközöket, amelyek között vannak hangvezérléses megoldások is, amellyel például az indexelés, ablaktörlő, világítás vezérelhető. Miután röviden ismertettük a legújabb innovációkat, amelyek támogathatják a mozgássérült sofőröket, rátérnénk arra, hogy a szakirodalom milyen előnyeit és hátrányait taglalja a mesterséges intelligencia használatának vezetés közben. Az átláthatóság és könnyebb olvashatóság érdekében ezt táblázatos formában közöljük, amelyet az 1. táblázat tartalmaz. A táblázatba a 2020-2025 között megjelent friss tanulmányokat gyűjtöttük össze.

Az elmúlt 5 év közleményei alapján a következő megállapításokat tehetjük: abban valamennyi közlemény egyetért, hogy nagy potenciál rejlik a mesterséges intelligencia bevonásában a közlekedésben, különösen a közlekedés szempontjából sérülékenyebb csoportok számára. Bár napjainkban is már nagyon sok kutatás-fejlesztés zajlik, azonban ezek többsége még kísérleti és kipróbálási fázisban tart. Nem minden rendszer elérhető kereskedelmi forgalomban. A járművek adaptációja szintén nehézségekbe ütközik (kézi vezérlők, pedál módosítások) hagyományosan mechanikusan vagy elektromechanikusan működnek és nem minden gyártó integrál „okos” MI-vezérlést ezekbe.

A következő korlátozó tényezőt az önvezető rendszerek szabályozási és jogi kérdései

jelentik. Ezek jogi háttere minden országban más-más szinten tart jelenleg, ami korlátozhatja a mozgássérültek számára nyújtható autonóm megoldások gyors elterjedését. Végül maguk a mozgássérült vezetők is sokszor bizonytalanok a mesterséges intelligencia vezetéssel támogató funkcióinak használatát illetően (Petrović és mtsai 2022; Bastola és mtsai 2024; Golbabaie és mtsai, 2024; Yousfi és mtsai, 2025).

Módszertan

A vizsgálat alapját egy 61 fős online kérdőíves kutatás képezte, amely mozgássérült személyek közlekedési tapasztalataira és biztonságérzetére fókuszált. A kérdőív az Oravecz (2024a, 2024b, 2025a, 2025b) által korábban kidolgozott felmérésekre épült. A válaszok feldolgozását SPSS szoftver segítségével végeztük, kereszttábla-elemzéseket alkalmazva a biztonságérzet és a vezetési gyakorlat kapcsolatának feltárására. A vezetési gyakoriság, a jogosítvány megszerzésének ideje és a járműhasználati szokások képezték az elemzés alapját (Patil és mtsai, 2024). Az alábbiakban a vizsgálat eredményeinek leíró-statisztikai szintű áttekintése következik.

Eredmények

Az eredmények alapján a megkérdezettek 89%-a rendszeresen vezet, és a többség (94%) biztonságban érzi magát sofőrként. Mindössze 10 fő számolt be arról, hogy utasként érzi magát biztonságban, és csupán 2 fő jelölte meg a tömegközlekedést mint legbiztonságosabb közlekedési módot. Érdekes, hogy az utasként önmagukat biztonságban érzők

1. táblázat. Az MI előnyei és hátrányai a mozgássérült vezetők számára a 2020-2025 megjelent közlemények alapján. Forrás: A szerző

Előnyök	Hátrányok / Kockázatok	Források
ADAS rendszerek növelik a reakcióképességet, csökkentik a baleseti kockázatot.	Függőség alakulhat ki, csökkenhet a figyelem.	Zeng, X., Wang, H., és Liu, Q. (2024)
MI-alapú akadályfelismerés növeli a helyzetfelismerés pontosságát.	Adatvédelmi és etikai aggályok.	Lee, S. és Park, J. (2023)
Autonóm/félaautonóm járművek mobilitást adnak súlyos mozgáskorlátozottság esetén.	Szenzorhibák és rendszermeghibásodások veszélyt jelenthetnek.	Martinez, A. és Gupta, R. (2022)
Személyre szabható járművezérlés a vezető fizikai állapotához igazítva.	Magas költségek és hozzáférési egyenlőtlenség.	O'Connor, L. és Chen, P. (2021)
MI-alapú akadálymentesítési funkciók támogatják a fogyatékossgal élők utazását az önvezető járművekben.	Kísérleti rendszerek, jogi és biztonsági bizonytalanságok.	Bastola és mtsai (2024)
Automatizált járművek növelhetik a mobilitást különféle fogyatékossgal élők számára.	Technológiai megbízhatósági kérdések és elfogadási nehézségek.	Dicianno és mtsai (2021)
Felhasználói igényekhez illeszkedő autonóm járműtervezés javíthatja a hozzáférhetőséget.	Különböző fogyatékossgai csoportok igényeinek komplexitása.	Dwyer és mtsai (2023)
Multimodális felületek segítik a látássérülteket az autonóm járművekből való biztonságos kiszállásban.	Technológiai túlterhelés, tanulási nehézségek.	Meinhardt és mtsai. (2025)
Önvezető járművek által kínált megoldások javítják a védett csoportok mobilitását.	Infrastruktúra korlátjai, működési problémák.	Zhong és mtsai (2024)
Automatizált járművek fenntartható mobilitási lehetőségeket biztosítanak fogyatékossgal élőknek.	Társadalmi elfogadottság, szabályozási kihívások.	Yousfi, Jacquet, és Métayer (2025)

fele szintén rendszeres vezető, ami arra utal, hogy a tömegközlekedési lehetőségek akadálymentesítésének hiánya miatt mégis inkább a vezetést választják. A több mint 5 éve jogosítvánnyal rendelkezők körében a biztonságérzet magasabb volt 86%, de még a kevesebb, mint 5 éve vezetők közül is 64%-a sofőrként érzi magát leginkább biztonságban.

Megbeszélés

Az adatok azt mutatják, hogy a vezetési gyakorlat és a szubjektív biztonságérzet között pozitív kapcsolat áll fenn. A biztonságérzetet nem csupán a járműtechnológiai tényezők, hanem a társadalmi és pszichológiai komponensek is alakítják. A mesterséges intelligencia alkalmazása a közlekedésben – például a vezetéstámogató rendszerek vagy az autonóm járművek révén – jelentős mértékben növelheti a mozgássérültek önállóságát. Ugyanakkor ezek a rendszerek új etikai dilemmákat is felvetnek, például a döntéshozatal társadalmi érzékenységevel kapcsolatban (Yousfi és mtsai, 2025; Booker és mtsai, 2024).

Következtetések

A vizsgálat rávilágít arra, hogy a mozgássérült vezetők biztonságérzetét erősen befolyásolja a vezetési rutin és az önálló közlekedés lehetősége. A mesterséges intelligencia alkalmazása a közlekedésbiztonságban kulcsszerepet játszhat a hozzáférhetőség és az inkluzívabb mobilitás elősegítésében. A jövőbeni fejlesztések során a technológiai megoldások mellett a pszichológiai és infrastruk-

turális tényezők integrált figyelembevétele szükséges. Fontos leszögeznünk azonban, hogy az MI nem helyettesíti, hanem kiegészíti az emberi felelősséget a közlekedésben – különösen a sebezhető csoportok esetében (Patil és mtsai 2024).

Irodalom

- Andrade, A., Escudero, M., Parker, J., Bartolucci, C., Seriani, S., & Aprigliano, V. (2024). Perceptions of people with reduced mobility regarding universal accessibility at bus stops: A pilot study in Santiago, Chile. *Case Studies on Transport Policy*, 16, 101190. Doi: <https://doi.org/10.1016/j.cstp.2024.101190>
- Bastola, A., Wang, H., Haeri Boroujeni, S. P., Brinkley, J., Moshayedi, A. J., & Razi, A. (2024). *Driving towards inclusion: A systematic review of AI-powered accessibility enhancements for people with disability in autonomous vehicles*. arXiv Preprint. Megnyitva: 2025.12.01. URL: <https://arxiv.org/abs/2401.14571>
- Booker, N., Holcombe, A., Gupta, B., Beard, G., Hitchings, J., Forrest, J., & Harris, E. (2024). The Safe System Approach and technology—what works? *Frontiers in Public Health*, 12, 1457616. Doi: <https://doi.org/10.3389/fpubh.2024.1457616>
- Dicianno, B. E., Sivakanthan, S., Sundaram, S. A., Satpute, S., Kulich, H., Powers, E., Deepak, N., Russell, R., Cooper, R., & Cooper, R. A. (2021). Systematic review: Automated vehicles and services for people with disabilities. *Neuroscience Letters*, 761, Article 136103. Doi: <https://doi.org/10.1016/j.neulet.2021.136103>

- Du, Z., Zhang, L., & Wei, H. (2023). Driver behavior analysis using AI-driven telematics: Emerging risk prediction models. *Transportation Research Interdisciplinary Perspectives*, 17, 100783.
- Dwyer, J., Gomez, R., Bubke, A., Cocks, K., & Paz, A. (2023). Designing Autonomous Vehicles for People with Disabilities: Surveying User Needs and Preferences. In *Proceedings of the 44th Australasian Transport Research Forum (ATRF)*.
- Golbabaei, F., Dwyer, J., Gomez, R., Peterson, A., Cocks, K., Bubke, A., & Paz, A. (2024). Enabling mobility and inclusion: Designing accessible autonomous vehicles for people with disabilities. *Transport Policy*. Preprint. Doi: <http://dx.doi.org/10.2139/ssrn.4784820>
- Gomez, R., Dwyer, J., Peterson, A. F., Bubke, A., Cocks, K., & Paz, A. (2023). On the road to enhancing transportation access for people with disabilities: A data review of accessible autonomous vehicles research. In *Proceedings of the 15th International Conference on Automotive User Interfaces and Interactive Vehicular Applications (AutomotiveUI '23)*. Doi: <https://doi.org/10.1145/3581961.3609867>
- International Transport Forum. (2022). *The Safe System Approach in Action*. OECD Publishing. Megnyitva: 2025.12.01. URL: <https://www.itf-oecd.org/sites/default/files/docs/safe-system-in-action.pdf>
- International Transport Forum. (2023). *Road Safety Annual Report 2023*. OECD Publishing. Megnyitva: 2025.12.01. URL: <https://www.itf-oecd.org/sites/default/files/docs/irtad-road-safety-annual-report-2023.pdf>
- Khan, M. N., & Das, S. (2024). Advancing traffic safety through the Safe System Approach: A systematic review. *Accident Analysis & Prevention*, 199, 107518. Doi: <https://doi.org/10.1016/j.aap.2024.107518>
- Knightley, A., Holcombe, A., Gupta, B., Beard, G., Hitchings, J., Forrest, J., & Harris, E. (2021). *PPR2049 – The impact of automated transport on disabled people: Summary report*. TRL & RiDC. Megnyitva: 2025.12.01. URL: https://www.trl.co.uk/uploads/trl/documents/PPR2049-Impact-of-Automation-on-Disabled-People_Summary-Report.pdf
- Közlekedésbiztonság. (2021, augusztus 16). *Mi is az a közlekedésbiztonság?* Megnyitva: 2025.12.01. URL: <https://vizsgakozpont.hu/hirek/mi-is-az-a-kozlekedesbiztonsag-20210816>
- Liu, L., Kar, A., Tokey, A. I., Le, H. T., & Miller, H. J. (2023). Disparities in public transit accessibility and usage by people with mobility disabilities. *Journal of Transport Geography*, 109, 103589. Doi: <https://doi.org/10.1016/j.jtrangeo.2023.103589>
- Lv, C., Zhao, Y., Zhang, J., & Na, X. (2023). Artificial intelligence for autonomous and assistive driving: A review of current technologies and future directions. *Transportation Research Part C: Emerging Technologies*, 157, 104285. Doi: <https://doi.org/10.1016/j.trc.2023.104285>
- Martinez, A., & Gupta, R. (2022). Reliability challenges in autonomous vehicle sensor technologies: A systematic review. *Transportation Technology and Safety*, 9(4), 221–239.

- Meinhardt, L.-M., Wilke, L., Elhaidary, M., von Abel, J., Fink, P., Rietzler, M., Colley, M., & Rukzio, E. (2025). *Light My Way: Developing and Exploring a Multimodal Interface to Assist People With Visual Impairments to Exit Highly Automated Vehicles*. arXiv. Megnyitva: 2025.12.01. URL: <https://arxiv.org/abs/2501.11801>
- Mobility in Motion. (2022). *Voice-controlled secondary driving functions for accessibility*. Megnyitva: 2025.12.01. URL: <https://www.mobilityinmotion.com/assistive-tech-voice-control>
- O'Connor, L., & Chen, P. (2021). Adaptive vehicle control systems for drivers with physical disabilities: Technological and social perspectives. *Assistive Mobility Review*, 6(2), 88–104.
- Oravecz, A. (2024a). Mozgássérült személyek jogosítvány szerzési és B kategóriás gépjármű vezetési tapasztalatai az érintettek és az oktatók szemszögéből. *OxIPO – Interdiszciplináris tudományos folyóirat*, 2024/2, 43–57. Doi: <https://www.doi.org/10.35405/OXIPO.2024.2.43>
- Oravecz, A. (2024b). Driving as a creative activity for disabled people. *OxIPO – Interdiszciplináris tudományos folyóirat*, 2024/4, 53–60. Doi: <https://www.doi.org/10.35405/OXIPO.2024.4.53>
- Oravecz, A. (2025a). A mozgáskorlátozott személyek gépjárművezetési tanulási folyamatának vizsgálata a résztvevők szubjektív jóllétének szemszögéből. *OxIPO – Interdiszciplináris tudományos folyóirat*, VII. évf. 2025/1, 25–33. Doi: <https://www.doi.org/10.35405/OXIPO.2025.1.25>
- Oravecz, A. (2025b). Mozgáskorlátozott személyek gépjárművezetési és tanulási folyamatának vizsgálata. *OxIPO – Interdiszciplináris tudományos folyóirat*, VII. évf. 2025/2, 89–96. Doi: <https://www.doi.org/10.35405/OXIPO.2025.2.89>
- Patil, D. S., Bailey, A., George, S., Ashok, L., & Ettema, D. (2024). Perceptions of safety during everyday travel shaping older adults' mobility in Bengaluru, India. *BMC Public Health*, 24(1). Doi: <https://doi.org/10.1186/s12889-024-19455-0>
- Petrović, Đ., Mijailović, R. M., & Pešić, D. (2022). Persons with physical disabilities and autonomous vehicles: The perspective of the driving status. *Transportation Research Part A: Policy and Practice*, 164, 98–110. Doi: <https://doi.org/10.1016/j.tra.2022.08.009>
- Pieres, C., et al. (2025). Comparative indicators for measuring road safety performance. *International Journal of Transport Safety*, 12(1), 1–18.
- Schmidt, T., & Weber, F. (2023). CogniSafe: An AI-enhanced driver monitoring system using deep learning and computer vision. *Proceedings of the AI & Mobility Conference*, 44–59.
- World Health Organization. (2023). *Global status report on road safety 2023*. Megnyitva: 2025.12.01. URL: <https://www.who.int/teams/social-determinants-of-health/safety-and-mobility/global-status-report-on-road-safety-2023>
- Yousfi, E., Jacquet, T., & Métayer, N. (2025). Automated vehicles and people living with a disability: Opportunities, challenges, and future directions for sustainable mobility. *Sustainability*, 17(13), 5941. Doi: <https://doi.org/10.3390/su17135941>

Zeng, X., Wang, H., & Liu, Q. (2024). AI-based driver assistance systems for people with physical disabilities: Challenges and opportunities. *IEEE Transactions on Intelligent Vehicles*, 9(2), 231–245. Doi: <https://doi.org/10.1109/TIV.2024.3351129>

Zhong, Y., Hu, L., & Ferreira, J. (2024). Autonomous vehicles and vulnerable road user mobility: Risks and opportunities. *Journal of Sustainable Transportation*, 18(1), 77–95.

DATA-DRIVEN TENSOR PRODUCT POLYTYPIC BASED MODEL FOR ENVIRONMENTAL NOISE MONITORING

Author(s) / Szerző(k):

Darwish Sultan
University of Debrecen (Hungary)

Zairi Housseem
University of Debrecen (Hungary)

Dénes Kocsis (Ph.D.)
University of Debrecen (Hungary)

E-mail:

housseem.zairi.hr6587@gmail.com

Cite: Sultan, Darwish; Housseem, Zairi and Kocsis, Dénes (2025): Data-driven
Idézés: Tensor Product Polytypic Based Model for Environmental Noise
Monitoring. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf.
2025/2. szám. 45-55.

Doi: <https://www.doi.org/10.35406/MI.2025.2.45>



This work is licensed under a Creative Commons Attribution-
NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0009

Reviewers: Anonymous reviewers / Anonim lektorok:

- Lektorok:**
1. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 2. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Abstract

This paper presents a compact matlab-based workflow for acquiring, organizing, and analyzing urban environmental noise data with a strict spatiotemporal structure. Multi-station sound pressure level (SPL) measurements are organized into a three-dimensional tensor, unfolded into a two-dimensional matrix, and analyzed using singular value decomposition (SVD) to extract dominant spatial and temporal patterns. Two lightweight normalization schemes cumulative normalized output (CNO) and separated normalized noise nodes (SNNN) are introduced to support low-rank reconstruction, anomaly detection, and decision-making. A case study conducted in Debrecen, Hungary, over an eight-week period (daily averages) demonstrates the effectiveness of the proposed methodology by revealing clear weekly patterns and midweek noise anomalies. The framework is scalable to dense sensor networks and can be integrated into smart city dashboards, offering a robust foundation for adaptive urban planning and real-time environmental noise management.

Keywords: Environmental noise; MATLAB; tensor data; singular value decomposition; smart city monitoring

Discipline: Environmental Engineering

Absztrakt**ADATVEZÉRELT TENZORSZORZAT-POLITÍPUS ALAPÚ MODELL
KÖRNYEZETI ZAJMONITOROZÁSHOZ**

Ez a tanulmány egy kompakt, MATLAB-alapú munkafolyamatot mutat be a városi környezeti zajadatok szigorú tér-időbeli struktúrában történő gyűjtésére, rendszerezésére és elemzésére. A több mérőállomáson végzett hangnyomásszint (SPL) mérések háromdimenziós tenzorba kerülnek szervezésre, majd kétdimenziós mátrixba kifejtve szingulárisérték-felbontással (SVD) kerülnek elemzésre a domináns térbeli és időbeli mintázatok feltárása érdekében. A tanulmány két könnyűsúlyú normalizálási eljárást vezet be – a kumulatív normalizált kimenetet (CNO) és a szétválasztott normalizált zajcsomópontokat (SNNN) –, amelyek támogatják az alacsony rangú rekonstrukciót, az anomáliadetektálást és a döntéstámogatást. Egy esettanulmány, amelyet Debrecenben végeztek egy nyolchetes időszak alatt (napi átlagértékekkel), igazolja a javasolt módszertan hatékonyságát azáltal, hogy jól elkülönülő heti mintázatok és a hét közepére jellemző zajanomáliákat tár fel. A keretrendszer skálázható sűrű szenzorhálózatokra, és integrálható okosváros-irányítópultokba, így robosztus alapot nyújt az adaptív várostervezéshez és a valós idejű környezeti zajkezeléshez.

Kulcsszavak: környezeti zaj; MATLAB; tenzoradatok; szingulárisérték-felbontás; okosváros-monitorozás

Diszciplína: környezetmérnöki tudomány

Environmental noise has become a pressing public health issue, particularly in densely populated urban areas where rapid population growth and infrastructure development contribute to rising noise levels. Extended exposure to high noise levels is associated with a variety of health problems, ranging from hearing loss to cardiovascular issues, as well as cognitive and sleep-related disturbances, making noise pollution a critical environmental and societal challenge (Basner et al., 2014; Sheikhmozafari et al., 2025; Chen et al., 2023). Traditional noise mapping approaches often rely on sparse, static measurements that capture only limited snapshots of the acoustic environment. Classical statistical techniques, such as linear or multiple regression models, are valued for their simplicity and low computational demands (Maulud & Abdulazeez, 2020). However, they struggle to capture the complex and dynamic spatiotemporal patterns of urban noise and tend to perform poorly when datasets are missing or irregular (Chen et al., 2015; Ahmed & Raihan, 2024).

The increasing deployment of real-time sensor networks and Internet of Things (IoT) based platforms has significantly improved the capacity for continuous acoustic monitoring and high-resolution data collection (Gunatilaka et al., 2025; Banaras & Nasir, 2025). While these systems generate valuable datasets, they also introduce challenges related to storage, processing, and interpretation. Addressing these challenges requires advanced analytical approaches that can efficiently model large-scale spatiotemporal

data while maintaining interpretability and scalability.

Recent advances in machine learning (ML) and artificial intelligence (AI) have been applied to environmental noise prediction and classification, often outperforming traditional statistical methods (Ali et al., 2022). Techniques such as neural networks and clustering algorithms provide powerful tools for uncovering hidden patterns, but they typically require extensive computational resources and large labeled datasets. Moreover, their “black-box” nature limits transparency, making it difficult for urban planners and policymakers to interpret and act on the results (Ali et al., 2022; Albaji et al., 2022; Szandala, 2023). A tensor-based framework offers an effective solution by organizing environmental data in a multidimensional form, allowing both spatial and temporal relationships to be retained. Spatiotemporal tensor decomposition methods, such as those based on Singular Value Decomposition (SVD), enable simultaneous data reduction, pattern extraction, and noise filtering, offering both computational efficiency and interpretability (Sudár et al., 2025; Auddy et al., 2024). Prior research on fuzzy modeling and rule-based reduction has also demonstrated the value of simplifying complex decision-making systems, which aligns closely with the goals of the present study (K & M. B, 2024; Jin et al., 2019; Dénes & Baranyi, 2000). Building on these concepts, this paper introduces a data-driven tensor product polytypic framework developed in MATLAB for urban noise monitoring and analysis. In the proposed

methodology, raw Sound Pressure Level (SPL) data collected from multiple monitoring stations are structured into a three-dimensional tensor. The tensor is then unfolded into a two-dimensional matrix and processed using SVD to reveal dominant spatiotemporal structures. This approach supports low-rank approximations, temporal modeling, and anomaly detection while remaining computationally efficient and interpretable. To validate the framework, a case study was conducted in Debrecen, Hungary. Over an eight-week period, the system successfully identified consistent weekly acoustic cycles and midweek anomalies with high precision. These results highlight the potential of the proposed framework as a decision-support tool for smart urban planning, adaptive noise regulation, and real-time environmental management systems.

Methodology

This section presents a systematic methodology for analyzing environmental noise using a 3D tensor-based approach in MATLAB. Earlier works concentrated on the straightforward reduction of heuristically generated rule bases.

The methodology aims to transform raw monitoring data into actionable insights by employing signal processing, mathematical modeling and visualization techniques.

Data Structure and Collection. Environmental noise levels were recorded using a network of fixed monitoring stations distributed across various urban zones. The collected data comprises daily average Sound Pressure

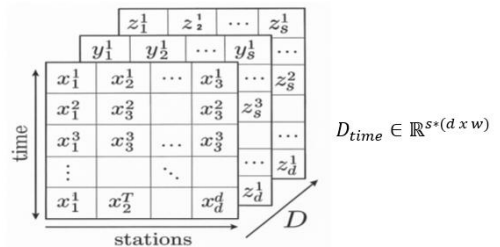
Levels (SPL) measured in decibels (dB), organized in a 3D tensor format to capture spatiotemporal dynamics:

- Spatial axis (s): Stations (e.g., $s = 12$)
- Temporal axis (d): Days of the week (Monday to Sunday, $d = 7$)
- Week index (w): Weekly observations (e.g., $w = 8$)

The tensor encapsulates the average SPL at each station (i), on each day (j), during each week (k).

Tensor Flattening and Matrix Construction. To facilitate decomposition techniques like SVD, the 3D tensor is unfolded into a 2D matrix: each row represents a monitoring station, and each column a specific day across all weeks. Example: the resulting tensor $\mathcal{D}$ has a shape of $(12 \times 7 \times 8)$, corresponding to 12 stations, 7 days per week, and 8 weeks. The tensor is then reshaped into a matrix $\mathbf{D}_{time}$ of shape (12×56) , preserving the spatial resolution while flattening the temporal dimensions for easier linear algebra processing (Figure 1).

Figure 1. Tensor-to-matrix unfolding for SVD analysis.



Singular Value Decomposition (SVD). SVD is applied to the flattened matrix:

$$D_{\text{time}} = U \cdot \Sigma \cdot V^T$$

Left singular vectors (spatial patterns)

Diagonal matrix of singular values

Right singular vectors (temporal patterns)

This factorization isolates the dominant modes of variation in noise across space and time.

Temporal Weight Function Extraction. Once the singular value decomposition (SVD) is performed on the unfolded noise matrix, the resulting right singular matrix contains vectors that represent the temporal modes of noise variation.

To analyze temporal behavior, the first two right singular vectors are selected. These vectors are transformed into weight functions:

$$\omega_1(S_a) = V_1(S_a)$$

$$\omega_2(S_a) = V_2(S_a)$$

Where $S_a \in [1, d \cdot w]$ represents a specific time index combining day and week. These temporal weight functions serve the following roles:

ω_1 Captures baseline trends in the noise profile, such as weekly periodicity and stable daily cycles.

ω_2 Captures secondary oscillations, reflecting more nuanced variations, such as day-to-day anomalies or mid-week disturbances.

In the case study conducted over May and June, these functions revealed:

ω_1 highlighted a repetitive weekday/week-end pattern, particularly peaking during workdays and dropping during weekends.

ω_2 emphasized periodic deviations, suggesting the presence of irregular local activities or transient noise sources (e.g., construction, traffic surges).

Together, these extracted functions provide the temporal structure necessary for downstream modeling and enable the decomposition of noise evolution into interpretable components.

Normalization Techniques. To enhance interpretability and robustness, two normalization techniques are integrated:

- Cumulative Normalized Output (CNO), which highlights long-term temporal trends and recurring daily patterns.
- Separated Normalized Noise Nodes (SNNN), which isolates short-term irregularities and high-intensity noise events.

Noise Pattern Modeling. The normalized weight functions are combined linearly to model the SPL signal: Where and control the contribution of each mode.

Model examples: Raw SVD; CNO; SNNN

Pattern Recognition and Decision Support. The extracted models support various practical applications:

- Recurring Pattern Detection: Identifying weekday-weekend cycles or seasonal effects.
- Anomaly Detection: Spotting unexpected events or deviations.
- Classification: Grouping days by similar noise behavior.
- Urban Noise Mitigation: Locating high-risk zones and optimal intervention times.
- Policy Evaluation: Measuring the impact of noise control strategies.

- Smart City Integration: Enabling real-time monitoring and alerts.

This methodology offers a comprehensive and scalable framework for environmental noise analysis using advanced signal decomposition in MATLAB.

Each model serves a specific purpose: trend extraction, anomaly detection, or signal reconstruction

Results And Discussion

Case Study: Station-Based Noise Analysis

To demonstrate the practical application of the proposed methodology, a case study was conducted using data from a single noise monitoring station situated in central Debrecen, Hungary. The location was chosen based on its suitable positioning and consistent recording performance.

The study covers the months of May and June (2025), totaling eight weeks. This duration enables identification of both transient acoustic disturbances and recurring weekly sound patterns. Data analyzed consisted of daily average sound pressure levels (SPLs), measured in A-weighted decibels (dB(A)), processed using the entire 3D noise analysis pipeline: tensor construction, matrix flattening, singular value decomposition (SVD), normalization, and temporal modeling.

This case study validates the methodology's ability to:

- Capture temporal dynamics
- Detect Anomalies
- Generate decision-supporting insights into noise management

Data Collection

The dataset originates from one Class 1 SV 307A noise monitoring station and covers 8 weeks (May-June), with a daily resolution. The structured format was a 3D tensor, reshaped into a 2D matrix for analysis (Table 1).

Table 1. Daily average SPL readings over 8 weeks [dB(A)]

Day	Week							
	W1	W2	W3	W4	W5	W6	W7	W8
Monday	54.9	53.3	52.0	52.3	53.2	55.4	56.9	56.9
Tuesday	54.5	53.2	53.5	51.4	54.6	56.4	53.4	56.2
Wednesday	55.1	52.9	54.6	53.6	55.3	55.4	55.5	53.1
Thursday	53.4	52.1	53.7	52.4	54.3	58.2	55.4	54.9
Friday	50.8	54.8	51.9	56.3	55.4	55.8	55.1	53.5
Saturday	50.4	55.2	51.6	52.3	53.6	51.7	52.9	52.5
Sunday	52.3	49.6	51.8	50.5	52.0	51.9	51.7	50.4

Results from MATLAB Simulation

Tensor Reshaping:

- Decomposition: SVD yields dominant vectors
- Normalization: Applied via CNO and SNNN methods

Model Outputs:

Raw model:

$$SP_{\text{mas}}(\text{sa}) = 233 \cdot w_1(\text{sa}) + 6 \cdot 10^{-14} \cdot w_2(\text{sa})$$

CNO model:

$$SP_{\text{mac}}(\text{sa}) = 44 \cdot w_1^{\text{cno}}(\text{sa}) + 60 \cdot w_2^{\text{cno}}(\text{sa})$$

SNNN model:

$$SP_{\text{snnn}}(\text{sa}) = \alpha \cdot w_1^{\text{snnn}}(\text{sa}) + \beta \cdot w_2^{\text{snnn}}(\text{sa})$$

Key Insights:

- Weekly cycles and anomalies confirmed
- CNO: Emphasizes long-term structure
- SNNN: Highlights daily variability

Table 2 shows model errors, and Figure 2 summarizes the weekly normalization comparisons.

Table 2. Model errors per week

Week	ER_snnn	ER_cno
1	$1.56e^{-13}$	$2.13e^{-14}$
...	...	...
8	$9.94e^{-14}$	$1.42e^{-14}$

Multi-Week Analysis

To better interpret temporal dynamics, average normalized curves (Figure 3) were computed for three periods: Weeks 1–4, Weeks 5–8, and Weeks 1–8. Both CNO and SNNN methods were used to construct representative sound profiles by linearly combining the first two temporal weight functions.

CNO curves highlighted smooth, long-term trends across days and weeks.

SNNN curves emphasized short-term variability and anomalies.

Key Patterns Identified:

- Midweek Peaks (Day 3 – Wednesday): Recurrent noise spikes linked to increased midweek activity.

- Weekend Calm (Days 6–7 – Saturday & Sunday): Notable decrease in noise, reflecting reduced urban activity.

These patterns remained consistent across all time segments, validating the reliability of both normalization approaches in capturing meaningful temporal structures in environmental noise.

Figure 2. Weekly normalization comparisons (CNO vs. SNNN)

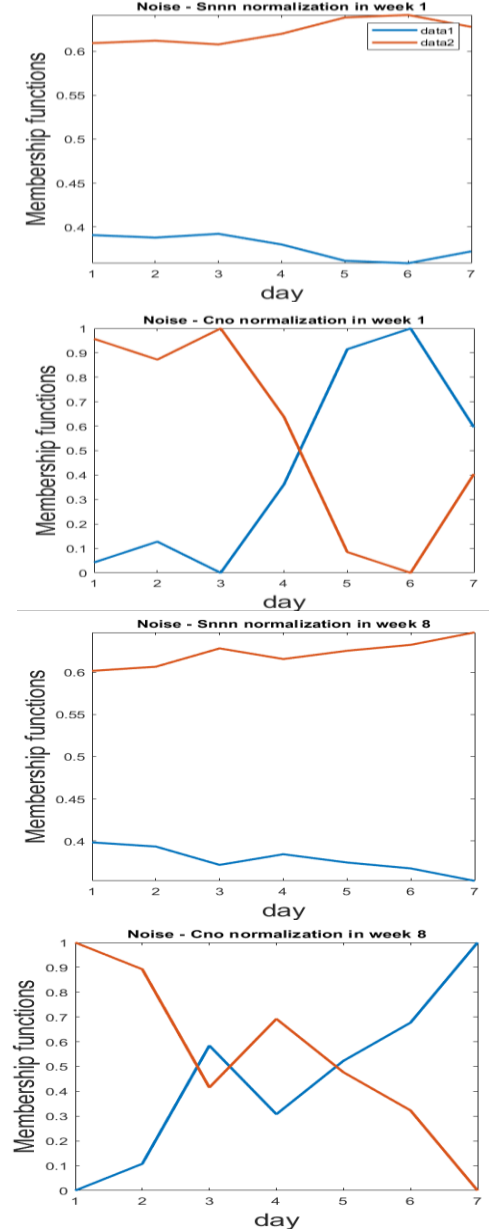


Figure 3. Average normalized values (W1–W8)

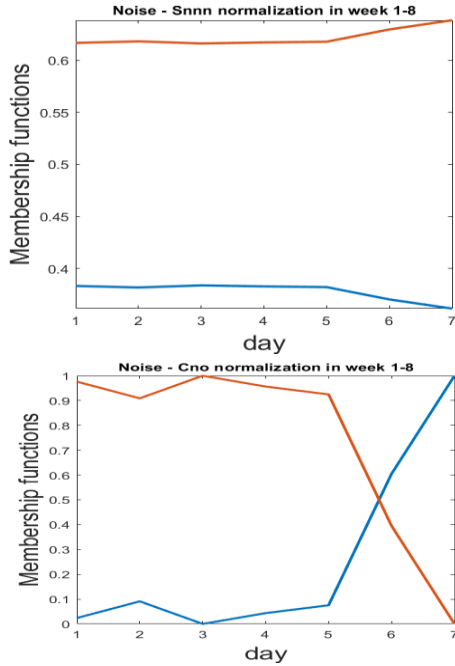
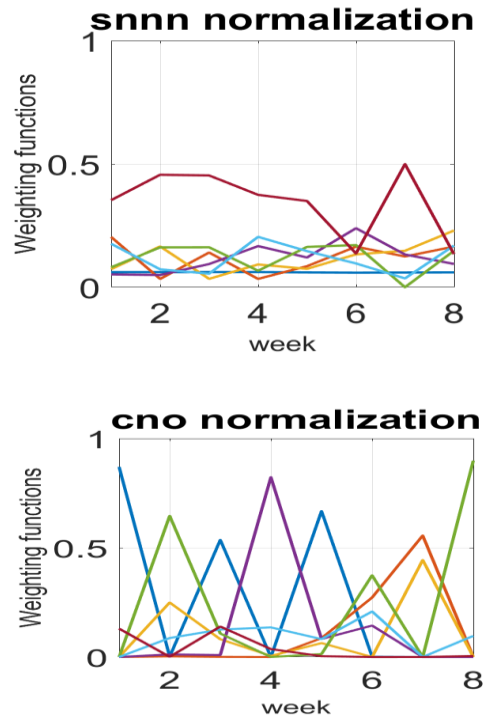


Figure 4. Normalized weighting functions



Weighting Function Identification

- SNNN: A dominant dark red curve models ambient noise trends best.
- CNO: The blue curve models peak midweek noise.

The normalized weighting functions can be seen in Figure 4.

Spatio-Temporal Surface Analysis

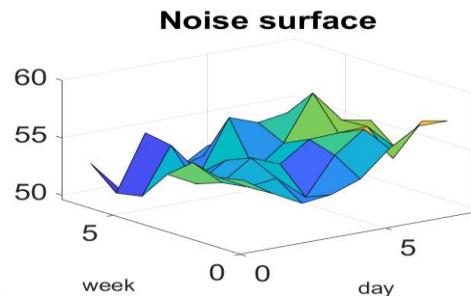
3D noise surface reveals (Figure 5):

- Midweek noise ridges
- Weekend valleys

SNNN: One dominant component

CNO: Chaotic peaks (local events)

Figure 5. Noise surface – 3D



To sum up, SNNN normalization provides a stable reference for ambient noise modeling, while CNO captures short-term noise disturbances (Table 3). These insights are essential for adaptive urban noise mitigation strategies.

Table 3. Summary of findings

Feature	SNNN	CNO
Pattern	Smooth, stable	Fragmented, volatile
Dominance	One major component	Multiple transient peaks
Usefulness	Long-term modeling	Local disturbance detection

Analysis and Interpretation

The MATLAB-based analysis of environmental noise data yielded several key insights into the temporal and structural behavior of urban sound patterns:

- **Dominance of a principal component:** The first singular value consistently captured more than 90% of the total variance, indicating the presence of a dominant and repetitive trend in the noise data. This component reflects the regular weekly rhythm of urban noise, with identifiable high and low cycles.
- **Secondary mode differentiation:** The second singular mode introduced moderate fluctuations, corresponding to behavioral transitions between weekdays and weekends. These shifts illustrate the natural alternation between high-activity (workdays) and low-activity (weekend) environments.
- **Temporal consistency and weekly structure:** The reconstructed SPmac surface model

demonstrated a stable weekly pattern across multiple weeks, with systematic noise peaks observed midweek (especially on day 3 or 4) and significant declines toward the weekend, confirming earlier statistical observations.

- **Detection of local noise anomalies:** The SNNN normalization model revealed distinct outlier contributions, associated with short-duration, high-intensity noise events. These include likely causes such as traffic congestion, public events, or construction, and were clearly identifiable due to the sharp concentration in the weighting and SC (spatial-temporal concentration) maps. Model stability and interpretability: While CNO normalization revealed highly chaotic and scattered contributions, SNNN normalization exhibited structured, dominant components — particularly one that aligned strongly with weekend silence. This confirms that silence modeling offers more stable and interpretable structures than raw noise amplitude.

These outcomes demonstrate the effectiveness of combining singular value decomposition (SVD), spatio-temporal modeling, and normalized weighting functions for extracting meaningful information from raw environmental data. The results are not only statistically robust, but also operationally relevant for practical noise monitoring and control systems.

Conclusion

The analysis of environmental noise data over eight consecutive weeks has highlighted the structured, cyclical nature of the examined

urban soundscape, with consistent patterns of midweek noise surges and weekend calm. Using advanced matrix decomposition techniques and membership function normalization (CNO and SNNN), the study successfully:

- Isolated week-specific anomalies and spikes, and
- Revealed a clear weekly acoustic rhythm.

Among the models, SNNN normalization proved the most effective in revealing structured, interpretable ambient noise patterns, while CNO normalization emphasized the irregular and reactive nature of noise events.

This methodological framework offers a valuable foundation for:

- Early-warning systems for noise pollution,
- Urban acoustic planning,
- Policy-driven interventions targeting specific temporal windows of noise vulnerability.

In sum, the work demonstrates that noise monitoring can go beyond raw decibel levels toward intelligent pattern recognition, strategic mitigation, and sustainable sound management in modern urban environments.

References

- Ahmed, I., & Raihan, A. S. (2024). *Spatiotemporal Data Analysis: A Review of Techniques, Applications, and Emerging Challenges* BT – *Multimodal and Tensor Data Analytics for Industrial Systems Improvement* (N. Gaw, P. M. Pardalos, & M. R. Gahrooei (eds.); pp. 125–166). Springer International Publishing. Doi: https://doi.org/10.1007/978-3-031-53092-0_7
- Albaji, A. O., Rashid, R. B. A., Sarijari, M. A. Bin, Salam, Z. Bin, Hamid, S. Z. A., & Ali, Y. H. (2022). *A Machine Learning for Environmental Noise Monitoring and Classification Using Matlab*. June 2024, 194–217. Doi: <https://doi.org/10.5281/zenodo.6640773>
- Ali, Y. H., Rashid, R. A., & Hamid, S. Z. A. (2022). A machine learning for environmental noise classification in smart cities. *Indonesian Journal of Electrical Engineering and Computer Science*, 25(3), 1777–1786. Doi: <https://doi.org/10.11591/ijeecs.v25.i3.pp1777-1786>
- Auddy, A., Xia, D., & Yuan, M. (2024). *Tensor Methods in High Dimensional Data Analysis: Opportunities and Challenges*. 1–38. Opened: 01.12.2025. URL: <http://arxiv.org/abs/2405.18412>
- Banaras, F., & Nasir, W. (2025). *Real-Time IoT-Based Monitoring of Mechanical Systems Using Acoustic and Vibration Data*. Opened: 01.12.2025. URL: <https://www.researchgate.net/publication/389338528>
- Basner, M., Babisch, W., Davis, A., Brink, M., Clark, C., Janssen, S., & Stansfeld, S. (2014). Auditory and non-auditory effects of noise on health. *The Lancet*, 383(9925), 1325–1332. Doi: [https://doi.org/10.1016/S0140-6736\(13\)61613-X](https://doi.org/10.1016/S0140-6736(13)61613-X)

- Chen, X. C., Faghmous, J. H., Khandelwal, A., & Kumar, V. (2015). Clustering dynamic spatio-temporal patterns in the presence of noise and missing data. *IJCAI International Joint Conference on Artificial Intelligence*, 2015-Janua(Ijcai), 2575–2581.
- Chen, X., Liu, M., Zuo, L., Wu, X., Chen, M., Li, X., An, T., Chen, L., Xu, W., Peng, S., Chen, H., Liang, X., & Hao, G. (2023). Environmental noise exposure and health outcomes: an umbrella review of systematic reviews and meta-analysis. *European Journal of Public Health*, 33(4), 725–731. Doi: <https://doi.org/10.1093/eurpub/ckad044>
- Dénes, I., & Baranyi, P. (2000). Inverse Model of DC-DC Converters. *IFAC Proceedings Volumes*, 33(28), 147–152. Doi: [https://doi.org/10.1016/s1474-6670\(17\)36825-8](https://doi.org/10.1016/s1474-6670(17)36825-8)
- Gunatilaka, D., Sawangphol, W., Charoenritthitham, T., Kanjanapoo, T., Burasotikul, T., & Pongprasit, K. (2025). An acoustic sensing system for noise monitoring and source identification using transfer learning. *Expert Systems With Applications*, 299(PB), 130014. Doi: <https://doi.org/10.1016/j.eswa.2025.130014>
- Jin, Y., Cao, W., Wu, M., & Yuan, Y. (2019). Accurate fuzzy predictive models through complexity reduction based on decision of needed fuzzy rules. *Neurocomputing*, 323, 344–351. Doi: <https://doi.org/10.1016/j.neucom.2018.10.010>
- K, N., & M. B, M. K. (2024). Fuzzy rule based classifier model for evidence based clinical decision support systems. *Intelligent Systems with Applications*, 22(May), 200393. Doi: <https://doi.org/10.1016/j.iswa.2024.200393>
- Maulud, D., & Abdulazeez, A. M. (2020). A Review on Linear Regression Comprehensive in Machine Learning. *Journal of Applied Science and Technology Trends*, 1(2 SE-Standard Journal Issues), 140–147. Doi: <https://doi.org/10.38094/jastt1457>
- Sheikhmozafari, M. J., Omari Shekaftik, S., Fasih Ramandi, F., Monazzam Esmacelpour, M. R., & Biganeh, J. (2025). A review of the studies investigating the effects of noise exposure on humans from 2017 to 2022: Trends and knowledge gaps. *Noise Mapping*, 12(1). Doi: <https://doi.org/10.1515/noise-2025-0015>
- Sudár, A., Berki, B., & Horváth, I. (2025). Fuzzy model-based analysis of user feedback for product development insights. *Helijon*, 11(12), e43537. Doi: <https://doi.org/10.1016/j.helijon.2025.e43537>
- Szandala, T. (2023). Unlocking the black box of CNNs: Visualising the decision-making process with PRISM. *Information Sciences*, 642(May), 119162. doi: <https://doi.org/10.1016/j.ins.2023.119162>

GONDOLATOK A MESTERSÉGES INTELLIGENCIÁRÓL TANÁROKNAK

Author(s) / Szerző(k):

Gyarmati G. Péter (Prof. emeritus, Dr.)
Stanford University (USA)
Simonyi Professor for the Public Understanding of Science and Professor
of Mathematics & Computer Science

E-mail:

gyarmati@gyarmati.dr.hu

Cite: Gyarmati G. Péter (2025): Gondolatok a mesterséges intelligenciáról tanároknak. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 57-64.

Idézés:

Doi: <https://www.doi.org/10.35406/MI.2025.2.57>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0010

Reviewers: Anonymous reviewers / Anonim lektorok:

Lektorok:

1. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
2. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

Az alábbiakban a mesterséges intelligencia forrásairól, várható fejlődéséről lesz szó. Kiinduló álláspontunk az emberi alkotás, amelyet a józan ész, a hatalom mindenhatóságra törekvése és a kapitalista gazdaság korlátlanága mozgat. A mesterséges intelligencia a tudomány mai állása szerinti antropomorf modell alapján építkezik, azaz utánózni igyekszik az emberi gondolkodást, viselkedést. A humán intelligencia bázisa az emberi agy, a mesterséges a számítógép. Az

emberek társadalmakba szerveződnek szabályok és az erkölcs szerint. A józan ész ellentmond a gépek ilyen önszervező képességgel való ellátásának és az emberi célok ezt nem is teszik szükségessé. A cikk igyekszik feltárni a két intelligencia különbözőségét, lehetséges konvergenciáját. Ezen közben mindig megállapítva az ember felsőbbrendűségét és felelősségét. Az etikátlant, a pusztítást, az ember – a hatalom – okozza, még ha intelligenciával is rendelkező gépekkel végzik. Ez az írás a MATE MAFIOK konferencián (Gyöngyös 2025. június 25-27.) elhangzott előadás anyaga.

Kulcsszavak: mesterséges intelligencia, gépi tanulás, humán intelligencia, formális logika, döntési automatika, veszély

Diszciplína: informatika

Abstract

THOUGHTS ON ARTIFICIAL INTELLIGENCE FOR TEACHERS

Below, we will discuss the sources and expected development of artificial intelligence. Our starting point is still human creation, which is driven by common sense, the striving for omnipotence of power, and the unlimited nature of the capitalist economy. Artificial intelligence is based on the anthropomorphic model according to the current state of science, i.e., it tries to imitate human thinking and behavior. The basis of human intelligence is the human brain, and the computer for artificial intelligence. People organize themselves into societies according to rules and morality. Common sense contradicts the need to provide machines with such self-organizing ability, and human goals do not require it. This article tries to explore the difference and possible convergence of the two intelligences. At the same time, it always establishes the superiority and responsibility of man. The unethical, the destruction, is caused by man, the power, even if it is done by machines with intelligence. This article is the material of a presentation given at the MATE MAFIOK conference (Gyöngyös, June 25-27, 2025.).

Keywords: artificial intelligence, machine learning, human intelligence, formal logic, automated decision-making, danger

Discipline: information technology

*Minden másképpen van,
mint ahogy tudni véljük:
...mert addig csűrítetek, begyezítetek,
Hasogatjátok, élesítitek,
Míg örültség vagy béklyó lesz belőle...*

(Karinthy Frigyes)

Korunk óriási divatja – minden média ezt hirdeti – a Mesterséges Intelligencia (MI) elnevezés. Az első mesterséges intelligencia kongresszust 1956-ban Dartmouthban tartották (v.ö.: McCarthy és tsai, 1955) és a LISP első változata már 1958-ban elkészült. A

remény hamar szertefoszlott, hamar kiderült, hogy az emberi okoskodás bármilyen leírása pusztán logika, sőt abban a pillanatban, hogy elkezd gépen működni algoritmussá válik, tehát nem intelligencia többé. Az algoritmus nem intelligencia, hanem állapotszekvencia, vagy rekurzor. Tipikus példa a perceptron, amely pusztán a neuron mesterséges modell közelítése, távol egy bizonyított definíciótól. Kétségbeesésre azonban nincsen semmi ok, mivel az eredmény – a perceptron modell és változatai – számos területen jól alkalmazható eszköz: felismerő és kereső algoritmusok stb.

Ma már senki sem tudja mit takar valójában ez a kifejezés. Megfigyelésünk szerint a világ legnagyobb számítástechnikai vállalkozóinak éppen kibocsátott alkotásaival találkozunk ezen a néven. A korábbi informatika címén tanított Office rendszert egyre inkább kiváltják újabb programalkotásokkal, melyek több szolgáltatást nyújtanak, korszerűbbek. Az informatikai oktatás ezek igényes, alapos alkalmazása. Nagyban hozzájárulnak ehhez az alkotó vállalkozók azzal, hogy úgynevezett oktatási változatot ingyen rendelkezésre bocsátanak. A tanuló és a generatív MI technikák legegyszerűbb változataival hasonló a helyzet és az oktatók fantáziájuktól függően dolgoztatják bennük hallgatóikat. A ma elért eredmények és a bennük rejlő lehetőségek jogosítanak minket a generatív mesterséges intelligencia elnevezés használatára. Az elnevezése és rövidítése: Artificial Generative Intelligence (AGI), magyarul Mesterséges Alkotó Intelligencia (MAI).

Ez a generatív MI még arra is alkalmas, hogy kiértékelje a munkát és verseny sorrendeket

hozzanak létre. Csúcsteljesítménynek számít, ha a feladatokkal a korlátokat keresik. Az előállító vállalatoknak – a hatalmas üzleti hasznon túl – éppen ez a másik célkitűzésük: kipróbálni, vizsgáztatni alkotásukat, hogy mennyire felelnek meg majd az üzleti követelményeknek. Így aztán születhetnek egymásra az újabb és újabb verziók. Az oktatók nem kis fáradtsággal ismerkednek az alkalmazásokkal és igyekeznek azt beilleszteni a tanrendbe, a hallgatók pedig kiképzett használói lesznek azoknak. Az oktatási eredmény alig mutat többre, mint a betanított munkási színvonal. Jó gondolat ez a multimilliomos software gyártó cégeknek, mert alkotásaik bevezetésekor már készen adott az emberi tényező.

Kérdések és válaszok

Tovább gondolva, ha egyáltalán lehetséges itt magasabb színvonalat elérni, az még mindig csak valamiféle utánzás lesz! Hol van az önállóság, az egyéni fantázia, ötlet, az igazi alkotó készség? Pedig a magasabb rendű oktatásnak ez lenne az elvárása. Az MI területén is sokkal inkább az alkotást kellene tanítani, mint például intelligencia algoritmusok, érzékelők szabályozó eljárásai, adatredukációs megoldások, valószínűségek, statisztika, eltérés minimalizálás, hiba felismerések, biztonság stb. Hasonló elméleti kérdések oktatása ad jövőt a diplomásoknak.

Tekintsük át röviden az oktatás legfőbb elemeit, amelyek minden korban befolyásolják és meghatározzák az elvárásokat:

1. A tudományos diszciplínák változása, fejlődése miatt a tantárgyak egyre bővülő tananyaggal néznek szembe. Néhány példa:

- nem-Euklideszi geometriák, terek, kvantum logika,
- hálózattudomány, kibernetika,
- kvantummechanika, részecskefizika,
- nanotechnológia, új anyagok.

2. A tudomány és technika aránya mindenkori problémát jelentett, ami egyidős az elmélet-gyakorlat vitával. A számítógép terjedésével azonosan egyre növekszik a technika aránya, ami a gépek és programjainak használatának tanítását jelenti. Megjelent a számítógéptechnika mellett a számítógéptudomány, a számítástechnika mellett pedig számítástudomány, valamint kibernetika néven a rendszer-, a szabályozás- és vezérlés elméletek technikai változata. Az informatika oktatás szinte kizárólag technikára korlátozódik, mert nagy és dinamikus ez az anyag is és nincs idő az elméleti megalapozásra. Megint az utánzásra, mások kiszolgálására kell hivatkoznom, amely sokkal inkább megfelel a közép-szintnek. Ennek teljesítéséhez megfelelő oktatók képzése elengedhetetlen.

3. A gőzgép felfedezése óta immár a negyedik ipari forradalomnak vagyunk részesei. Általában a gőz, az olaj, az elektromosság, és végül a számítógép jellemzi az egyes ipari forradalmakat. Korunkat digitális kornak is nevezzük. Az alapvetően új alkotások életünk minden területére betörnek. Néhány ilyen technika és módszer a hálózatok, adatbázisok, robotok, önvezető eszközök, nanotechnológia, műanyagok, kvantummechanika, energiatárolás, virtualitás, biotechnológia stb. A

technológiák olyan gyorsan újulnak meg, amellyel a biológiai ember már szinte lépést sem tud tartani. Az emberiség szétszakad, lesznek ezek művelői, kevesen alkotói, és lesznek leszakadók. Jó lenne idejekorán felismerni ki hová tartozik – ki, miben tehetséges –, mert az „egyformaság” kizárólag a színvonal csökkentésével lehetséges. Végre ismét mindenkit tehetségesnek tartunk. A feladat annak megtalálása, felismerése, hogy miben. Tehetőség-gondozási feladat a megtalált tehetség kellő irányba terelése és az alkalmas képzés. A magyarság mindig büszke volt kiválóságaira és azok alkotásaira. Az egyetemek alapkövetelménye legyen a kiválóság, oktatókban, oktatásban, elért eredményekben egyaránt.

4. Az extenzív eszközeink kimerültek. Az oktatásra szánt idő már nem növelhető tovább. Intenzív módként a szakterületek szétválasztása adódik, mint lehetőség. Jó ideje ismert a STEM megoldás, valamint a generációk sajátosságainak feltárása. A STEM a Science, Technics, Engineering, Mathematics rövidítése. Véleményem szerint, tömören fogalmazva, Science: aki figyeli a természetet, Technics: aki megvalósít, Engineering: aki kitalál, Mathematics: aki elképzel – mindegyik más tehetséget feltételez. A Z-generáció a XXI. század elejének szülöttei, akik már a digitális korszakban élnek, annak technikáján nőttek fel. A STEM a különböző tehetségek feltárásában, majd a hozzáillő szakterületre való terelésében rejlik. Az így kiválasztottakhoz már sokkal közelebb fog állni az illető terület legújabb eredményeinek megismerése és a hajlam az aktív alkotásra.

5. A megoldáshoz a tehetség idejekorán való felismerése szükséges, vagyis ehhez vissza kell térnünk a középiskolába, mert ez a kor a tehetség kibontakozása. Felkészült tanárookra van szükség, akik képzése fontos főiskolai, egyetemi feladat. Sajnos, de tudomásul kell vennünk, hogy ez többéves feladat, és a következő generációra lesz hatásos. És akkor, addig csak a betanított munka – hogyan használjuk mások alkotását – marad?

Múltból a jelenbe

A múlt század közepéig a tudományos felfogás szerint a természet formálta a világot:

*„Megpillantjuk a tudomány szemével,
Bölcsőjében mint forgott a világ,
S hogy a természet fortélyos kezével
A nagy káoszt miként formálta át.”*

(Spencer)

Az ipari forradalmak ezt a szemléletet fokozatosan átalakították és egyre inkább az ember alkotta technikai fejlődés vált a tudományos világ meghatározójává. Már 1954-ben Szentgyörgyi (1988) egy konferencián feltette a kérdést: képes lesz-e a biológiai ember utolérni a maga alkotta technikai világát, mi lesz, ha mind-örökre lemarad? És ez a lemaradás azóta is egyre növekszik. Lag theory (lemaradás elv): mindig van lemaradás egy rendszer bevezetése során az új rendszer és az alkalmazói között, amely azonban rövidesen kiegyenlítődik. Ha a fejlődés gyorsabb, mint a kiegyenlítés, akkor a lemaradás véglegessé válik. A lemaradást ma már generációkban mérjük.

Az utóbbi évek hatalmas emberi újításai ezt a folyamatot igazolja, a technika egyeduralmódóvá válik és a tudomány az igényeinek a kiszolgálója lesz. Ez az ember alkotta technika mára betört az ember legbensőbb tudati világába a gondolkodás, az intelligencia, a döntés, a kivitelezés irányítása modellezésével. Ez a technika a McCulloh és Pitts által alkotott neuron-modellen alapul és a Neumann (1965) által kigondolt számítógépen realizálódik. Működése a Wiener által megfogalmazott kibernetikán alapul a Turing definiálta állapotokon keresztül, amelyek mindig lehetségesek, de csak bizonyíthatatlanul végesek. Alapszabály McCulloh-Pitts és Neumann felvetése, Markov normalizációs tétele, a Turing tézis, a Church tézis. Tehát, ha bármi, amit elegendően pontosan és egyértelműen megfogalmazunk, akkor az számítógépen realizálható. A feladatot tehát leképeztük az emberi nyelvre, mert azzal fogalmazunk, tehát az „elegendő pontosság és az egyértelműség” nyelvi követelmény, egész pontosan egy olyan nyelv, vagy nyelvi átmenet, amelyet az ember és a gép egyaránt megért. Ugyanakkor tudjuk, hogy a nyelv az emberek közötti kommunikáció eszköze és pontatlansága, egyedisége, többértelműsége miatt a félreértések okozta tragédiák világa is.

Akkor hogyan alkothatunk mégis használható eljárásokat? Az ember találékony és kitalálta a kísérletezés, próbálkozás technikáját, ehhez nem feltétel a pontosság és az egyértelműség. Az elért eredményeket összevethetjük a kívánttal és a döntést most még az emberre bizzuk. Módosítás, javítás és a követő újabb próbálkozás – új verzió – a fejlődési

folyamat. A mesterséges intelligencia, a gépi tanulás a maguk területén ezt a lehetőséget tartalmazzák. Sőt a generatív, az alkotó rendszereket az alkalmazó is taníthatja, formálhatja saját elképzeléseihez. Az eredményesség egyelőre kétséges, mindenesetre anyagilag az alkotók számára felettébb hasznos. De hát „kísérleti fázisban” vagyunk!?

Összefoglalás

Összefoglalva jól láthatjuk, hogy az AGI-t előállítók érdekei jól érvényesülnek, jelentős eredményeket sikerült elérni az alakfelismerésben, beszédértésben, a videó és audio technikákban, adatkezelésben stb. Soha nem látott méretű feldolgozási és tároló kapacitások – felhő – állnak rendelkezésre és kínálják az AGI algoritmusok ezreit a szoftvergyártók a felhasználók számára legtöbbször előfizetéssel – ez az alkotók milliárdos haszna – és miként már említettem, mindezek oktatási változatai is elterjednek.

A megoldások nincsenek hiba nélkül, az eredmények hozzáértéssel, figyelemmel, megfelelő szintű ellenőrzéssel alkalmazhatók csak, mégis óriási a hajlam az automatikus döntés bevezetésére, emberi beavatkozás nélküli alkalmazására számos területen.

Különösen a hatalom részesei és a média teszik ezt és egyre gyakoribb ezáltal a fake news, a zero tolerancia és társai. Olyan szintet képesek ezzel elérni, hogy a valóság kiderítése lehetetlenné válik és a mérlegelés, a megértés az emberi tulajdonságok kizáródnak. A MI technika az emberek többsége fölé nő, uralkodik felettük, holott a demokrácia mellé-

rendelő felépítésű. Kimondhatjuk, hogy nem a technikával, hanem az egyes emberek gyakorlatával van a baj.

A modernebb technikák bevezetése – régóta tudjuk – emberi munkát tesz feleslegessé, a valóságban átalakítja azt könnyebb, kényelmesebb tevékenységgé, egyre inkább igényelve a felnőttképzést, ami szintén újdonsága a közép- és felső szintű oktatásnak.

A fejlődést nyomon követő képzés alapvető fontosságú és minden késlekedés generációk lemaradását, az egész társadalom leépülését okozhatja.

Zárómegjegyzés

Az előadásban elhangzottakat igazolandó megemlítem, hogy az USA-beli és a Brit tevékenységem – még a hivatalos időn túl is – a tehetség felfedezéséről és segítéséről szólt. Munkaidőben tanszékvezetőként, illetve kutatási igazgatóként vezettem a tanárok továbbképzési programjának kialakítását. Szabadidőmben a magyar hagyományokra épülő Math Circle, illetve Math Cafe programok kialakítását és működését szerveztem. Nemzetközileg is elismert magyar szervezetek, kiadványok a SICC régi és újabb kiadásai, valamint a ma is élő KÖMÁL füzetek; az Arany Dániel és Rácz László versenyek adták az ötletet (1. ábra).

A MATH CIRCLE a Stanford University tanárai által létesített „vasárnapi iskola” a matematikus gondolkodáshoz iskolásoknak. A MATH CAFE a Cambridge University NRIC termésettudományi kutatócsoportjának kezdeményezése.

1. ábra: A matematikai tehetséggondozás és az MI néhány szellemi műhelye. Forrás: a szerző



Prof. Dr. Gyarmati a Stanford University emeritus professzora

Mindezek célja az akkori Z generációból a matematika és informatika terén tehetséges iskolások segítése volt. Mára ezek és más hasonló programok országos méretűek az USA-ban és Angliában.

Végül e cikk témája iránt éreklődők figyelmébe ajánlhatók még a következő szakirodalmak: Challoner (2000), Gyarmati (2019, 2022, 2023a,b,c), Gyarmati és Mező (2024), Pólya (1963), illetve webodalak:

<https://gyarmati.eu>

<https://nrich.maths.org>

<https://www.superhuman.ai/c/confirmation>

<https://hspublishing.org/JRECS/article/view/133>

<http://www.sciencepg.org/article/10.11648/j.ijis.20221105.11>

<https://hspublishing.org/JRECS/article/view/168>

<https://hspublishing.org/JRECS/article/view/283>

Irodalom

- McCarthy, J., Minsky, M., Rochester, N., Shannon, C. (1955): *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence* (1955). Opened: 1.12.2025. URL: <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>.
- Neumann, J. Von (1965): *Selected Papers*. KJK, Budapest. 114-115.
- Szentgyörgyi A. (1988): *Selected Papers*. Gondolat, Budapest.
- Challoner, J. (2000): *The Brain*. MacMillan, New York.
- Gyarmati P. (2019): Gondolatok a mesterséges intelligencia, gépi tanulás kapcsán. *Mesterséges intelligencia – interdiszciplináris folyóirat*, I. évf. 2019/1. szám. 31–39. Doi: <https://www.doi.org/10.35406/MI.2019.1.31>
- Gyarmati P. (2022): Gondolatok a mesterséges intelligencia, a gépi tanulás kapcsán II. rész. *Mesterséges intelligencia – interdiszciplináris folyóirat*, IV. évf. 2022/2. szám. 9-25. Doi: <https://www.doi.org/10.35406/MI.2022.2.9>
- Gyarmati P. (2023a): Gondolatok a mesterséges intelligencia, a gépi tanulás kapcsán (III. rész). *Mesterséges Intelligencia – interdiszciplináris folyóirat*, V. évf. 2023/2. szám. 25-38. Doi: <https://www.doi.org/10.35406/MI.2023.2.25>
- Gyarmati P. (2023b): Húsvéti gondolatok a Mesterséges Intelligenciáról: egy keresztényi közelítés. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, V. évf. 2023/2. szám. 19-24. Doi: <https://www.doi.org/10.35406/MI.2023.2.19>
- Gyarmati P.G. (2023c): *Szilánkok*. TCC Computer Studio, Wien.
- Gyarmati P. és Mező F. (2024): Egy kísérlet a generatív mesterséges intelligenciát reprezentáló technológiával. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VI. évf. 2024/1. szám. 47-58. Doi: <https://www.doi.org/10.35406/MI.2024.1.47>
- Pólya, G. (1963): Mathematical Methods in Natural Sciences. MAA, Stanford. *Stanford Encyclopedia of Philosophy: Artificial Intelligence*. Opened: 1.12.2025. URL: <https://plato.stanford.edu/archives/fall2022/entries/artificial-intelligence/>

**MESTERSÉGES INTELLIGENCIA ÉS KREATIVITÁS:
HATÁROK, KÜLÖNBSÉGEK ÉS ÚJ FOGALMI KERETEK**

Author(s) / Szerző(k):

Csajbók Gábor
Eszterházy Károly Katolikus Egyetem
(Magyarország)

E-mail:

csajbok.gabi@gmail.com

Cite: Csajbók Gábor (2025): Mesterséges intelligencia és kreativitás: határok,
Idézés: különbségek és új fogalmi keretek. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 65-74.
Doi: <https://www.doi.org/10.35406/MI.2025.2.65>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0011

Reviewers: Public Reviewers / Nyilvános Lektorok:
Lektorok: 1. Erzsébet Lestyán (Ph.D.), Gál Ferenc Egyetem
2. Olteanu Lucian (Ph.D.), Gál Ferenc Egyetem (Magyarország)

Anonymous reviewers / Anonim lektorok:
3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Abstract

ARTIFICIAL INTELLIGENCE AND CREATIVITY: BOUNDARIES, DIFFERENCES, AND NEW CONCEPTUAL FRAMEWORKS

This paper raises the question of whether artificial intelligence can be considered creative. This is also difficult to determine due to the many different definitions of creativity. The standard definition of creativity – which focuses on novelty and value – does not provide a sufficiently precise definition. Therefore, we need to supplement it with new criteria, including intentionality and authenticity. This paper presents the fundamental differences between human and artificial systems, such as differences in cognitive abilities and functionality, and structural and functional differences between biological and artificial "nervous systems". We point out that artificial intelligence cannot be considered creative in the same way as humans, and therefore, new, more precise conceptual frameworks are needed. Due to this difference, human and artificial creative abilities are not competitors, but rather complementary phenomena.

Keywords: creativity, AI

Discipline: psychology, information technology

Absztrakt

E tanulmány azt a kérdést veti fel, hogy a mesterséges intelligencia kreatívnak tekinthető-e. Ennek meghatározása a számos, egymástól eltérő kreativitás-definíció miatt is nehéz. A standard kreativitás definíció – amely az újdonságra és az értékre koncentrál – nem nyújt kellően pontos meghatározást. Ezért új kritériumokkal, többek között az intencionalitással és a hitelességgel szükséges kiegészítenünk azt. Jelen tanulmány bemutatja az emberi és mesterséges rendszerek közötti alapvető különbségeket, például a kognitív képességek és a funkcionalitás eltéréseit, a biológiai és mesterséges „idegrendszerek” strukturális és működésbeli eltéréseit. Rámutatunk, hogy a mesterséges intelligencia nem tekinthető az emberrel azonos módon kreatívnak, ezért új, pontosabb fogalmi keretekre van szükség. Ezen eltérés miatt az emberi és a mesterséges kreatív képességek nem versenytársak, hanem egymás kiegészítésére alkalmas tulajdonságok.

Kulcsszavak: kreativitás, MI

Diszciplína: pszichológia, informatika

A kreativitásra az emberi élet számos területén szükség van (Desdevises, 2025, Mező, 2017), éppen ezért fontos lehet, hogy milyen eszközökkel lehet az emberi kreatív folya-

matokat támogatni, és hogy a mesterséges intelligencia (MI) fejlődése felülmúlhatja-e ezt az emberi képességet. A MI létrejöttével és fejlődésével újra felmerült a kérdés, hogy nem

csak az ember lehet kreatív, valamint fontos kutatási téma lett a mesterséges intelligencia potenciáljának és határainak vizsgálata (lásd: Desdevises, 2025; Mező, 2024). Habár kreativitásfogalmunk mindig is antropocentrikus volt, a kreativitást meghatározó emberi képességnek tekintve, a MI megjelenése előtt is voltak vizsgálatok, melyek azt kutatták, vajon csak az ember képes-e ilyen folyamatokra. Ezen vizsgálatok során az emberi és állati képességeket hasonlították össze, vizsgálva, hogy ezek miben térnek el egy-mástól (Da Pelo, 2025). A MI megjelenése bonyolultabbá tette a kérdést, hogy vajon csak az ember kreatív-e, hiszen míg az állatok csak egyszerűbb feladatokra találtak megoldást, a mesterséges intelligencia már képes létrehozni művészeti alkotásokat is, melyek olykor nehezen megkülönböztethetők az emberi alkotásoktól (Da Pelo, 2025). Az utóbbi években egyre több figyelmet kapott ez a téma, mely meghatározó lehet abban, ahogyan ezeket az eszközöket használjuk.

E cikk azt vizsgálja, hogy a MI mennyire felel meg a kreativitás elméleti követelményeinek, általános meghatározásainak, elemzi annak működési mechanizmusait, végül pedig arra a következtetésre jut, hogy a MI nem mint versenytárs, hanem mint kiegészítő, alternatív módon kreatív partner értelmezhető.

A kreativitás meghatározásának nehézsége

Habár a kreativitás egy elterjedt fogalom, valójában mégis nagyon nehéz pontosan megfogalmazni, hogy pontosan mit is értünk alatta

(Ismayilzada et al., 2024), s komoly konceptualizációt és operacionalizációt érintő problémákkal küzdő fogalomról van szó (Mező és Mező, 2021, 2022, 2023). Számos kutató rámutatott arra, hogy a kreativitás mint pszichológiai fogalom mindmáig ellenáll az egyértelmű definíciónak és a világos operacionalizálásnak, így nem létezik széles körben elfogadott, egységes meghatározása. Ennek ellenére a szakirodalomban számos definíciós kísérlet született, és több próbálkozás történt ezek rendszerezésére is (Mező, 2017; Mező és Mező, 2011). Jó példája ennek Aleinikov és mtsai (2000) könyve is, amelyben a szerző több mint 100 definíciót sorol fel a kreativitás meghatározására (Ismayilzada et al., 2024). Az, hogy a kutatók és gondolkodók ilyen eltérő módon definiálják a kreativitást azt eredményezi, hogy az arról való tudományos gondolkodás is nagyon különböző lehet (Mező, 2017).

A standard definíció korlátai és kibővítése

A filozófusok és pszichológusok körében találhatóunk egy széles körben elfogadott, ún. „standard definíciót” (Ismayilzada et al., 2024). E meghatározás szerint a kreativitás az a képesség, amely újszerűséget (novelty) és értéket (value, hasznosság/megfelelőség) hoz létre (Ismayilzada et al., 2024; Magurean et al., 2024; Desdevises, 2025; Schubert, 2021; O’Toole & Horvát, 2024; Mező, 2017). Ha ennél a fogalomnál maradunk, akkor azt mondhatjuk, hogy a MI is kreatív, hiszen az általa létrehozott termékek gyakran bizonyul-

nak újszerűnek és hasznosnak (Runco, 2023). Ez az állítás azonban csak akkor állja meg a helyét, ha nem vizsgáljuk meg pontosan azt a folyamatot, melyben létrejöttek – ez ugyanis sokban eltér az emberi alkotófolyamatoktól (Runco, 2023)

Ennek oka az, hogy a standard kreativitás-fogalom önmagában félrevezető lehet, hiszen csak a létrehozott termékre koncentrálnak, és figyelmen kívül hagyja, hogy ez a termék hogyan jött létre (Runco, 2023). Ha ebből a meghatározásból indulunk ki, akkor akár bizonyos geológiai vagy kémiai folyamatokat is kreatívnek kellene tartanunk, hiszen általuk új formációk jönnek létre, amelyek az ember számára hasznosak lehetnek (például: arany, oxigén). Két választási lehetőségünk van: a) vagy megmaradunk a standard fogalomnál, de így olyan folyamatokat is kreatívnek nevezzük, melyek valójában sokban különböznek az emberi folyamattól, vagy pedig b) pontosítjuk ezt a meghatározást, és így jobban láthatjuk a különbségeket (Runco, 2023).

A kreativitás pontos meghatározásához tovább kell mennünk és nem csak azt kell néznünk, hogy „mi” jön létre (product), hanem azt is, hogy „hogyan” (process). Ha ezt a mechanizmust nem vizsgáljuk, akkor egy pontatlan képet kapunk (Runco, 2023).

A kreativitásfogalom egyéb kritériumai

Amint az láthattuk, a kreativitás pontos meghatározásához további kritériumok szükségesek az újszerűsége és a hasznosságon túl (1. táblázat).

1. táblázat. Az ember és a MI összehasonlítása néhány kreativitás kritérium alapján.

Kreativitás kritérium	Ember	MI
Újdonság	✓	✓
Értékesség/Hasznosság	✓	✓
Intencionalitás	✓	✗
Autentikusság	✓	✗
Belső Motiváció	✓	✗
Választás Lehetősége	✓	✗
Spontaneitás	✓	✗
Önazonosság (Self)	✓	✗

Ahhoz, hogy ez a képesség/folyamat emberi értelemben legyen értelmezhető, ki kell egészítenünk egyéb szempontokkal is a kreativitásra vonatkozó koncepciókat. Ebben segít minket Schubert, 2021 és Runco, 2023 cikke.

Intencionalitás.

Schubert (2021) szerint négy komponens szükséges ahhoz, hogy kreativitásról beszéljünk: az alkotás képessége, az intencionalitás, a kontextus, valamint az új és hasznos termék létrejötte. Az emberi kreativitásban az intencionalitás kulcsfontosságú: az emberi alkotó tudatosan dönt arról, hogy létrehoz valamit, és belső motivációk vezérlik. Ezzel szemben a napjainkban ismert mesterséges intelligencia rendszerei algoritmusok és emberi parancsok alapján működnek, így kérdéses, mennyiben tekinthetők intencionálisnak (Ismayilzada et al., 2024). Runco (2023) tovább árnyalja ezt a képet, kiemelve, hogy az intencionalitás magában foglalja a belső motivációt, a választás

lehetőségét és a spontaneitást – olyan tényezőket, amelyek az emberi kreatív folyamatok lényegi részei, de a MI esetében hiányoznak.

Hitelesség

Runco (2023) egy másik kritériumot is felvet: a hitelességet (authenticity, itt önzonosság). Az emberi kreativitásban a hitelesség az önzonosság, az eredetiség és a valódiság kifejeződése. Maslow és Rogers szerint az autentikus egyén önfelfogadó, nem manipulálja túlzottan saját kifejezését, és képes szabadon megnyilvánulni. A MI esetében azonban nincs „self”, amelyből fakadhatna az önzonosság, ezért bár képes lehet szűrők (filter) nélküli kifejezésre, ez nem azonos azzal, amit pszichológiai értelemben hitelességnek nevezünk (Runco, 2023).

A mesterséges és emberi kreativitás közötti alapvető különbségek

Ahhoz, hogy még inkább lássuk az ember és a mesterséges intelligencia határait, hasznos, ha áttekintjük az ezek alapjául szolgáló mechanizmusokat és a közöttük fennálló különbségeket.

A kognitív hiány és a funkcionalitás

A mesterséges intelligencia rendszerei – például a nagy nyelvi modellek – nem tudatosan működnek, vagyis nincs saját öntudatuk, nincsenek szándékaik vagy önálló megértésük, mint az embernek. Ha öntudattal bírnának, akkor feltételeznünk kellene, hogy

rendelkeznek valamiféle intencionalitással (szándékossággal), sőt akár hiedelmekkel, vágyakkal, félelmekkel vagy tudattal. Mivel azonban a MI viselkedését nem belső szándékok, hanem előre meghatározott algoritmusok és mintafelismerés irányítják, ezek a rendszerek funkcionálisak, nem pedig strukturálisan tudatosak. Más szóval: az MI nem „akar” semmit – csak végrehajtja a betanult mintákat (Da Pelo, 2025). A MI által létrehozott alkotásokból hiányzik a tudatosság (consciousness), érzelmi mélység, személyes tapasztalatok, autonómia és egyéni intencionalitás (Da Pelo, 2025; Linares-Pellicer et al., 2025).

Luciano Floridi szerint (2019, 2023) a MI igazi újítása nem abban rejlik, hogy intelligenssé vált, hanem abban, hogy képes elválasztani bizonyos feladatokat az emberi Intelligenciától, vagyis el tud végezni olyan műveleteket, amelyekhez korábban intelligencia kellett, habár „ő” maga mégsem rendelkezik vele (Da Pelo, 2025). Bár a MI működése hasonlít az emberi kreativitás modelljeire, a mesterséges intelligencia mégsem tekinthető valóban kreatívnak a szó emberi értelmében, mivel hiányzik belőle a tudatos kognitív feldolgozás, az intencionalitás és a szubjektív tapasztalat.

A mesterséges (ANN)

és a biológiai (BNN) idegrendszer

A mesterséges intelligencia fejlődésében két fő irányzat alakult ki: a gépközpontú és az ember által inspirált megközelítés. A mai nagy nyelvi modellek (Large Language Model, LLM), mint a ChatGPT vagy a Google Gemini az első kategóriába tartoznak. Működésük alapját az emberi agy által inspirált mes-

terséges neurális hálózat (Artificial Neural Network, ANN) adja, de ezek nem rendelkeznek valódi szerkezeti vagy funkcionális hasonlósággal a biológiai idegrendszerrel (Biological Neural Network, BNN), amely elektrokémiai mechanizmusokra épül (Da Pelo, 2025; Linares-Pellicer et al., 2025). Az ANN-ok matematikai műveleteken és numerikus optimalizáláson alapulnak, szemben az emberi agy elektrokémiai és tapasztalat-alapú tanulási

folyamataival (2. táblázat). Ezért a mesterséges intelligencia kreativitása statisztikai minták kombinációján alapul. Ugyanakkor az emberi megértés, erkölcs és kontextus-ismeret hiányzik a mesterséges intelligenciából, még akkor is, ha a feldolgozó képességének többszörösének megfelelő adatbázisokkal képes dolgozni (v.ö.: Linares-Pellicer et al., 2025, Ismayilzada et al., 2024).

2. táblázat. Biológiai (BNN) és mesterséges (ANN) neurális hálózatok összehasonlítása. Forrás: Linares-Pellicer et al. (2025, p. 4).

Jellemző	Biológiai neurális hálózat (BNN)	Mesterséges neurális hálózat (ANN)
Alapegység	Biológiai neuron (sokféle típus, összetett struktúra)	Mesterséges neuron/perceptron (matematikai funkció)
Jel típusa	Elektrokémiai impulzusok	Numerikus értékek
Tanulási mechanizmus	Szinaptikus plaszticitás (Hebb-féle tanulás), strukturális változások, neurogenesis	Paraméterek/súlyok módosítása
Tudás tárolása	A szinapszisok erősségében és a hálózat szerkezetében elosztva	Numerikus paraméterekben (súlyok, biasok) elosztva
Információ forrása	Szenzoros bemenetek, testbe ágyazott tapasztalat, egész életen át tartó interakciók	Tanító adathalmazok (internetes szövegek/ képek)
Környezet	Testbe ágyazott, helyhez kötött, biológiai, evolúciós	Testetlen, matematikai, algoritmikus
Működési alapelv	Biológiai/elektrokémiai folyamatok	Matematikai függvények, statisztika

Megj.: BNN: Biological Neural Network, ANN: Artificial Neural Network

Boden kreativitás-típusai és a transzformatív határ

Nemcsak a kreativitás definíciójából kiindulva, hanem a kreatív folyamat (process) oldaláról megközelítve is felfedezhetünk különbségeket az ember és a mesterséges intelligencia között. Margaret A. Boden (1998) cikkében három típusát különbözteti meg a kreativitásnak:

1. Kombinatorikus kreativitás (combinatorial creativity): meglévő ötletek, fogalmak, elemek új módon történő kombinálása. Azt mondhatjuk, hogy a MI képes erre, amikor például képeket vagy zenéket hoz létre korábbi minták alapján. Ez az a szint, amelyben a MI ma a legerősebb, felülmúlva az emberi képességeket, hiszen hatalmas adatbázisokkal dolgozik.

2. Feltáró kreativitás (exploratory creativity): létező lehetőségek terének felfedezése, vagyis új variációk vagy megoldások keresése egy adott rendszeren belül. Bizonyos értelemben a MI is képes erre, amikor különféle kombinációkat fedez fel meglévő minták alapján.

3. Transzformatív kreativitás (transformational creativity): a legmagasabb szint, amelyben a szabályrendszer vagy a fogalmi tér változik meg, és így olyan ötletek, művek születnek, amelyek korábban elképzelhetetlenek voltak (Magurean et al., 2024).

A mai MI modellek képesek imitálni az emberi kreativitást, és annak kombinatorikus és feltáró típusát meg is tudják valósítani, azonban egyelőre nem képesek a transzformatív kreativitásra, amikor is nem a meglévő elemekből, hanem azoktól teljesen füg-

getlenül jön létre egy produktum (Magurean et al., 2024; Linares-Pellicer et al., 2025, Runco, 2023).

Co-creativity:

a mesterséges intelligencia mint partner

Bár a MI működési folyamata sok szempontból különbözik az emberi kreativitástól, nem egyszerűen csak utánzása, hanem alternatívája annak (Linares-Pellicer et al., 2025): a MI működési módja a matematikán és a mintázatok felismerésén, és nem a biológián és a személyes tapasztalaton alapszik. Ez a nézőpont ahhoz a felismeréshez vezet, hogy a két működési mód nem kizárja egymást, hanem egymás kiegészítője is lehet.

Számos tanulmány kimutatta, hogy a MI nagy segítség lehet kreatív folyamatokban alkalmazva. Különösen is alkalmas egyes sok időt igénylő feladatok teljesítésére. Ezáltal a művész, tudós, alkotó jobban tud fókuszálni egyéb fontos folyamatokra, amíg a MI elvégzi a technikai részt. Emellett a mesterséges intelligencia új ötletekkel segítheti a kreatív folyamatokat, új nézőpontokkal kiegészítve azt (Magurean et al., 2024).

„Ez a perspektíva, amelyet a folyamatban lévő diskurzus szempontjából alapvetőnek tartunk, a generatív mesterséges intelligencia rendszereket nem független feltalálókként, hanem a kollektív emberi kreativitás és tudás kifinomult feldolgozóiként tekinti – olyan rendszerekként, amelyek összefoglalják és új formákká alakítják közös digitális örökségünket” (Linares-Pellicer et al., 2025, 8.). Ez a megközelítés eltolja a hangsúlyt az eszköz

(tool) felől a kollaborátor felé, bevezetve a “co-creative” folyamatot (O’Toole & Horvát, 2024). Ahelyett, hogy lecserélnék vagy felülmúlnák az emberi kreativitást, ezek az eszközök a legjobban arra használhatók, hogy támogassák és fejlesszék az ötletalkotási folyamatot (Desdevises, 2025).

Az alkotás folyamata erőteljesen meghatározó a létrehozott termékben, hiszen ennek eredményeként jön létre az új alkotás. Ezért a lényegileg különböző folyamatok lényegileg különböző terméket hoznak létre, még akkor is, ha ezek külsőleg nagyon hasonlítanak, vagy olykor akár nehezen megkülönböztethetők egymástól (Runco, 2023). Felmerülhet a dilemma, hogy amikor egy új, kreatív tartott terméket hozunk létre, elég-e, ha csak a kész outputot tartjuk fontosnak, vagy azt is, hogy ez hogyan jött létre, van-e mögötte személyes tapasztalat, érzelmek, belső motivációk. (Runco, 2023; Mező, 2015; Mező és Mező, 2019).

Bár a MI a transzformatív kreativitás emberi értelemben vett formáját nem valósítja meg, alternatív működése miatt főként olyan kreatív folyamatokban válhat értékessé, ahol a kombinatorikus és feltárási kreativitás különösen hasznos.

Új megfogalmazások

Az eddig áttekintett érvek alapján a mesterséges intelligenciát nevezhetjük alternatív módon kreatívnek, mert bár vannak dolgok, amelyben hasonlít, sok szempontból különbözik az emberi kreativitástól – például a személyes tapasztalat, az öntudatot és egyéni intenciók aspektusából (Linares-Pellicer et al.,

2025). Amint a mesterséges intelligencia elnevezés nagyon elterjedt, érdemes a mesterséges kreativitás, és hasonló, pontosabb megfogalmazást nyújtó elnevezések használata is, mint pszeudó-kreativitás, hogy pontosan írjuk le a MI-t (Runco, 2023). Ezek a megfogalmazások kifejezik a két képesség közötti különbséget, sőt arra is utalnak, hogy ezek, jól alkalmazva, nem egymást kizáró, vagy egymással versengő tulajdonságok, hanem képesek kiegészíteni, segíteni egymást.

Összefoglalás

A tanulmány alapkérdése az volt, hogy vajon kreatívnek nevezhető-e a mesterséges intelligencia. Mint arról szó volt, ennek eldöntése amiatt is nehéz feladat, mert a kreativitás meghatározására számos eltérő definíció létezik. A standard definícióból kiindulva szükséges volt, hogy azt egyéb szempontokkal is kiegészítsük, mert önmagában nem ad mélyreható meghatározást. Az újdonság és érték (hasznosság) alapvető kritériumokat kiegészítettük az intencionalitással és a autentikussággal, hogy átfogóbb képet kapjunk. A MI és az ember közötti főbb különbségek áttekintésekor áramutattunk a kognitív hiány és a funkcionalitás, valamint a mesterséges és biológiai idegrendszer eltérő jellemzőire. Ezt követően Boden (1998) felosztását vizsgáltuk meg, amely alapján a mesterséges intelligencia nem képes a transzformatív kreativitásra. A tanulmány fő megállapítása, hogy a MI nem nevezhető az emberi kreativitással azonos módon kreatívnek, de a human és a mesterséges kreativitás nem feltétlenül vetélytársa egymásnak, hanem egymást kiegészíthetik.

Irodalom

- Aleinikov, A.G., Kackmeister, S. and Koenig, R. (2000). *Creating Creativity: 101 Definitions (what Webster Never Told You)*. Alden B. Dow Creativity Center Press. Megnyitva: 2025.12.01. URL: <https://books.google.ch/books?id=M2QpAAAACAAJ>
- Boden, M.A. (1998). Creativity and artificial intelligence. *Artificial Intelligence*, 1998, Volume 103, Issues 1–2, August 1998, Pages 347-356. Doi: [https://doi.org/10.1016/S0004-3702\(98\)00055-1](https://doi.org/10.1016/S0004-3702(98)00055-1)
- Da Pelo, M. (2025). Artificial creativity: Can there be creativity without cognition? *AI & Society*. Advance online publication. Doi: <https://doi.org/10.1007/s00146-025-02682-3>
- Desdevises, J. (2025). The paradox of creativity in generative AI: High performance, human-like bias, and limited differential evaluation. *Frontiers in Psychology*, 16, Article 1628486. Doi: <https://doi.org/10.3389/fpsyg.2025.1628486>
- Floridi, L. (2019). What the near future of artificial intelligence could be. *Philos Technol* 32(1):1–15. Doi: <https://www.doi.org/10.1007/s13347-019-00345-y>
- Floridi L. (2023) The ethics of artificial intelligence: principles, challenges, and opportunities. Oxford University Press
- Ismayilzada, M., Paul, D., Bosselut, A., & van der Plas, L. (2024). *Creativity in AI: Progresses and challenges* (arXiv preprint arXiv:2410.17218). Doi: <https://doi.org/10.48550/arXiv.2410.17218>
- Linares-Pellicer, J., Izquierdo-Domenech, J., Ferri-Mollà, I., & Aliaga-Torro, C. (2025, April 10). *We Are All Creators: Generative AI, Collective Knowledge, and the Path Towards Human-AI Synergy* (Preprint). arXiv. Doi: <https://doi.org/10.48550/arXiv.2504.07936>
- Magurean, I. D., Brata, A. M., Ismaiel, A., Barsan, M., Czako, Z., Pop, C., Muresan, L., Jurje, A. I., Dumitrascu, D. I., Brata, V. D., Leucuta, D. C., Stanculete, M., & Popa, S. L. (2024). Artificial intelligence as a substitute for human creativity. *Journal of Research in Philosophy and History*, 7(1), Article 7. Doi: <https://doi.org/10.22158/jrph.v7n1p7>
- Mező F. (2024): Módszertani útmutató mesterséges intelligenciák szöveggeneráló teljesítményének összehasonlító vizsgálatára fókuszáló tehetséggondozó programok számára. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VI. évf. 2024/2. szám. 63-72. Doi: <https://www.doi.org/10.35406/MI.2024.2.63>
- Mező F. és Mező K. (2011): *Kreatív és iskolába jár!* K+F Stúdió Kft., Debrecen.
- Mező F. és Mező K. (2019): Interdiszciplináris kapcsolódási lehetőségek a mesterséges intelligenciára irányuló cél-, eszköz- és hatásorientált kutatáshoz. *Mesterséges intelligencia – interdiszciplináris folyóirat*, I. évf. 2019/1. szám. 9–29. Doi: <https://www.doi.org/10.35406/MI.2019.1.9>

- Mező F. és Mező K. (2023): A személyre jellemző flexibilitás pontozásának háromfokú skálát alkalmazó módszere. *OxIPO – interdiszciplináris tudományos folyóirat*, 2023/1. 9-24. Doi: <https://www.doi.org/10.35405/OXIPO.2023.1.9>
- Mező K. (2015): *Kreativitás és élménypedagógia*. Kocka Kör, Debrecen.
- Mező K. (2017). *A kreativitás időbeli aspektusai*. Doktori disszertáció. Debreceni Egyetem, Humán Tudományok Doktori Iskola, Debrecen.
- Mező K. és Mező F. (2021). *A Körök- és a Szokatlan használat teszt magyar értékelő táblázatainak revideációja*. K+F Stúdió Kft., Debrecen
- Mező K. és Mező F. (2022). A hazai kreativitáskutatás trendjei, főbb vizsgálati kérdései. *Alkalmazott Pszichológia*, 22(2), 21-34. Doi: <https://www.doi.org/10.17627/ALKPSZ.ICH.2022.2.21>
- O'Toole, K., & Horvát, E. Á. (2024). Extending human creativity with AI. *Journal of Creativity*, 34(2), Article 100080. <https://doi.org/10.1016/j.joc.2024.100080>
- Runco, M. A. (2023). AI can only produce artificial creativity. *Journal of Intelligence*, 11(2), 22. Doi: <https://doi.org/10.1016/j.jint.2023.100022>
- Schubert, E. (2021). Creativity is optimal novelty and maximal positive affect: A new definition based on the spreading activation model. *Frontiers in Neuroscience*, 15, Article 612379. Doi: <https://doi.org/10.3389/fnins.2021.612379>

MÓDSZERTANI TANULMÁNYOK

METHODICAL STUDIES

WHEN ARTIFICIAL INTELLIGENCE GENERATED CONTENT BECOMES WEAPON-GRADE COMMUNICATION

Author(s) / Szerző(k):

Dávid Horváth
Ludovika University of Public Service
(Hungary)

E-mail:

david@netokracia.hu

Cite: Horváth, Dávid (2025): When Artificial Intelligence Generated
Idézés: Content Becomes Weapon-Grade Communication. *Mesterséges
Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 77-98.
Doi: <https://www.doi.org/10.35406/MI.2025.2.77>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0012

Reviewers: Public Reviewers / Nyilvános Lektorok:

- Lektorok:**
1. Jakusné Harnos Éva (Ph.D.), Ludovika University of Public Service (Hungary)
 2. László Teknős (Ph.D.), Ludovika University of Public Service (Hungary)

Anonymous reviewers / Anonim lektorok:

3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Abstract

This methodological paper explains how AI-supported, multimodal content can escalate into weapon-grade communication. After outlining the limits of detection-centric responses, it proposes an WGC-matrix that operationalizes credibility hijacking and cognitive overload at the intersection of intent, channel, and cognitive impact. The framework is validated through 12 case studies spanning text, image, audio, and video modalities.

Keywords: weapon-grade communication (WGC), generative AI, credibility hijacking, cognitive overload, hybrid warfare

Disciplines: military science; communication and media studies; criminology; security studies

Absztrakt

AMIKOR A MESTERSÉGES INTELLIGENCIA ÁLTAL GENERÁLT TARTALOM FEGYVERNEK MINŐSÜLŐ KOMMUNIKÁCIÓ

A tanulmány módszertani választ ad arra, miként válik a mesterséges intelligencia által generált multimodális tartalom fegyvernek minősülő kommunikációvá (FMK). A detekció-központú megközelítések korlátai után egy FMK-mátrixot javasolunk, amely a szándék–csatorna–kognitív hatás metszetében operacionalizálja a hitelesség-eltérítést és a kognitív túlterhelést. A keret 12 esettanulmányon kerül validálásra.

Kulcsszavak: fegyvernek minősülő kommunikáció (FMK), generatív mesterséges intelligencia, hitelesség-eltérítés, kognitív túlterhelés, hibrid hadviselés

Diszciplínák: hadtudomány; kommunikáció- és médiatudomány; kriminológia; biztonságpolitika

Content generated by artificial intelligence (AI) is forcing security policy and social science analysis onto a new path because, while the logic of influence is often familiar, the scale and pace of implementation are taking a qualitative leap. It is not just that more messages appear faster: the entire information ecosystem (platform architecture, recommendation systems, moderation, media routines, user psychology) begins to function as an operational space. In this sense, weapon-grade communication cannot be identified by

pinpointing "individual lies," but rather by the fact that content production, the perception of authenticity, and the structural conditions for dissemination are shifting together.

The "Firehose of Falsehood" in the age of algorithms: from quantity to quality

According to the classic view of agenda-setting (AS), the media space not only reflects reality, but also thematizes it: it does not tell us what to think, but what to think about

(McCombs, Shaw, 1972). Since the 2010s, this has been happening in a platform environment where the agenda is also organized by algorithmic ranking systems and feedback mechanisms. The "firehose of falsehood" (FoF) propaganda model describes this as high-volume, multi-channel, rapid and repetitive content distribution, often aimed not at persuasion but at overload and disruption of the decision-making environment (Paul, Matthews, 2016). The question is therefore not "how much" false content there is, but when quantity turns into qualitative impact: thematization, emotional attunement, and ultimately the erosion of public trust.

At the same time, platform logic also indicates that mere scaling is not automatically successful: dissemination takes place through the filters of recommendation systems and network structures, so organic reach is not simply a function of the number of posts produced. It can be empirically measured that the network of connections and ranking simultaneously limits and shapes encounters with differing views (Bakshy, Messing, Adamic, 2015). This is where the methodological dilemma becomes clear: traditional frameworks often stick to direct statements and "persuasion," while AI-supported FoF logic often plays on the long-term destruction of common points of reference.

In the military science community, this is organically linked to the doctrinal idea of shaping the information environment (SIE): the goal is not to "make a single message true," but to shape the context in which the recipient is forced to interpret it (UK Ministry of

Defence, 2013). The limitation of the diagnosis is therefore that traditional tools typically respond to content-related statements (fact-checking, source analysis, labeling), while the impact at the WGC level often manifests itself in the structure of the agenda, the rhythm of attention, and the erosion of trust.

When we can't even believe our own eyes: the crisis of multimedia credibility

In a multimodal disinformation environment, credibility is rarely determined by a single medium: text, images, audio, and video can all contribute to building a sense of authenticity. One of the central consequences of information disorder is that the boundaries between deliberate deception, misleading context, and manipulated content become blurred, while reactions are rapid and corrections are delayed (Wardle, Derakhshan, 2017). In this sense, credibility hijacking (CH) is not merely technical manipulation, but a communicative operation: the evidential nature (photo/video) lends borrowed authority to narratives that would otherwise be more easily dismissed as text.

The deepfake phenomenon (DF) encapsulates this shift: synthetic media not only spreads false claims, but also makes the status of evidence debatable, which is both a risk to democratic decision-making and a challenge to national security (Chesney, Citron, 2019). The blind spot is therefore twofold: the logic of "detectable fakes" is too narrow in a fast, first-impression-optimized dissemination cycle; and the crisis of credibility consists not

only in "believing a fake video," but also in the recipient later turning to the evidence-based media with general skepticism. In other words, the erosion of trust can still occur even if the specific fake is subsequently exposed.

Blind spots in current research models: why does misinformation slip through the filter?

The research blind spot often arises when we perceive the phenomenon as a short-term attitude change and underestimate the long-term recalibration of the cognitive environment. According to the illusory truth effect (ITE), mere repetition—even with existing knowledge—can increase credibility because processing fluency activates "familiarity = truth" shortcuts (Fazio, Brashier, Payne, Marsh, 2015). If we combine this with FoF tempo and platform rhythm, it becomes clear that AI does not produce false statements, but rather variants and repetitions that work together to lower the credibility threshold.

This is accompanied by the source monitoring (SM) thesis: memory stores not only content but also source cues, but these fade over time, leaving "source-free content" (Johnson, Hashtroudi, Lindsay, 1993). AI-assisted production scales this very phenomenon: the same narrative returns on new platforms, in new languages, and in new styles, while it becomes increasingly difficult to reconstruct where the audience first encountered it. Traditional diagnostic tools are often structurally delayed: by the time the refutation is ready, the source references have already evaporated from the attention cycle.

From this point of view, AI "slip through the filter" is not only a technical detection issue, but also a methodological gap: current models often do not measure consistently enough (i) the ecosystem effect resulting from scaling, (ii) the chains of multimodal authentication, and (iii) the long-term cognitive consequences. This is where the WGC framework comes in: it does not promise "yet another taxonomy," but rather the ability to address the phenomenon simultaneously at the levels of communication theory, criminology, and military science—with the same set of questions, in a comparable manner, with falsifiable statements.

Based on the above diagnosis, AI-assisted influence is not simply "more disinformation," but a shift in the operating parameters of the information ecosystem: scaling, multimodal authentication, and delayed correction together produce the WGC-level effect. It follows that the conceptual framework of a single discipline is not sufficient to grasp the phenomenon: communication theory agenda and attention mechanisms, criminological operational patterns, and military SIE logic must be read within a common framework. The following chapter therefore provides a targeted literature review of the conceptual and methodological approaches used in international and domestic research to address this problem of "more disinformation" — and identifies the blind spots that need to be covered by the subsequent analysis matrix.

Literature map: from propaganda to generative warfare (overview of publications)

Artificial intelligence-supported, multi-modal forms of weapon-grade communication (WGC) can be understood if we place the phenomenon not as an isolated technological novelty, but within the historical arc of propaganda, disinformation, and hybrid warfare research. In this arc, generative AI does not rewrite the rules of the game from scratch, but accelerates and extends the mechanisms already described: content production, the illusion of authenticity, and the scaling of dissemination (McCombs, Shaw, 1972; Paul, Matthews, 2016; Horváth, 2023a). The chapter organizes the literature review around three focal points: (i) the connections between classical schools of disinformation and hybrid warfare; (ii) the trap of technological determinism; (iii) the credibility crisis surrounding multimodality and "evidential media."

Classical schools of disinformation and the connection between them and hybrid warfare

Classic models of propaganda (Lasswell, 1927, 1948; Bernays, 1928; Jowett, O'Donnell, 2019) described communication in terms of a goal-channel-effect logic, where the source of influence is mostly institutionally identifiable. In this tradition, disinformation is an operational tool of propaganda: it deliberately distorts or falsifies information to reduce cohesion, weaken morale, and disrupt decision-making (Rid, 2020).

Cold War "active measures" (Kalugin, 1994; Barron, 1983; Andrew, Mitrokhin, 2000) indicated early on that communication is not merely a side effect, but an independent operational dimension. Fake documents, front organizations, and long-term narrative building created an information environment in which the recipient's "reality" could be gradually shifted.

21st-century hybrid warfare literature captures this relationship at the ecosystem level: military and non-military tools, overt and covert operations, and the toolkits of state and non-state actors form a mutually reinforcing system (Hoffman, 2007; Kramer, Speranza, 2017; NATO, 2016). In this framework, fake news is not a matter of press ethics, but a tactical tool that targets perception and decision-making processes (Horváth, 2022; 2023a; 2023b).

Generative AI acts as a qualitative accelerator at three points in this historical arc. (1) It radically reduces the cost and time of content production, making large-volume, multi-channel logics (FoF) easier, more sustainable, and more adaptable to operate (Paul, Matthews, 2016). (2) It reinforces the illusion of authenticity: illusory truth effect (Fazio et al., 2015) and source monitoring errors (Johnson, Hashtroudi, Lindsay, 1993) scale in multimodal, automated environments. (3) Communication in the hybrid space can increasingly be interpreted as a "weapon": it not only accompanies kinetic operations, but often triggers or prepares them (Horváth, 2024).

Subchapter synthesis (short): the classic propaganda–disinformation–hybrid warfare arc provides a stable conceptual basis for the WGC framework; the novelty of generative AI is not its intent, but its scale, speed, and automation of multimodal authentication.

*The trap of technological determinism
in current research*

Discourses surrounding generative AI often slip into technological determinism: as if the tool itself explains disinformation, the crisis of trust, or political polarization. Classic critiques (Winner, 1980; Morozov, 2013) point out that technology operates in a socio-technical space, between business, political, and institutional incentives.

Platform research (Gillespie, 2018; Couldry, Hepp, 2017) also emphasizes that algorithmic ranking and moderation build invisible rules: they amplify certain content and hide others. Generative AI enters this environment: its training data corpora, product strategies, and regulatory gaps together shape the information ecosystem in which WGC-type tactics operate.

Cybersecurity-focused approaches often treat the issue as a detection problem (e.g., watermarking, AI detectors; Kirchenbauer et al., 2023), while some social science analyses treat AI as a homogeneous threat, making little distinction between military, criminological, and commercial logics. The trap of determinism is therefore twofold: (i) it overestimates the "omnipotence" of the tool, (ii) it underestimates the strategic actors and the operational configuration.

This study addresses this from the perspective of the WGC framework and the WGC matrix: it does not ask "what AI can do," but rather who uses it as a weapon, with what intentions, with what set of channels, and under what ecosystem conditions (Horváth, 2023c; Horváth, 2024). The emphasis thus shifts from content tagging and detection to the exploration of structures relevant from a security policy perspective.

Subchapter synthesis (short): the risk of generative AI is not due to the "magic power" of the tool, but to the combination of platform logic + actor incentives + operational goals; the WGC matrix positions itself to analyze this socio-technical configuration.

*Multimodality and evidentiary media:
research directions in credibility hijacking*

For a long time, "evidential media" (photos, audio, video) were the pillars of credibility in the press, law, and political communication. Research in visual culture (Mitchell, 1994; Hariman, Lucaites, 2007) indicated early on that visual evidence is always a framed and interpreted sign, while meme culture (Shifman, 2014) demonstrated the logic of rapid recontextualization.

The legal and national security debate surrounding deepfakes and synthetic media is based on these insights: deepfakes pose a threat to privacy, democracy, and national security (Chesney, Citron, 2019). Furthermore, well-known elements of cognitive mechanisms are reinforced in multimodal environments: repetition can increase credi-

bility (Fazio et al., 2015, while source cues become worn out (Johnson, Hashtroudi, Lindsay, 1993).

Wardle and Derakhshan (2017) consistently treat different types of manipulated visual content within their framework of “information disorder”: recontextualization, false captions, modified and generated images. The novelty of generative AI is that the “evidential” content is no longer necessarily a reframing of existing material, but often a document illusion created from scratch.

In my own WGC framework, I have called this “credibility hijacking”: a tactic that uses existing trust structures (platforms, media, visual conventions) to elevate manipulated content to the status of evidence (Horváth, 2024). This leads to the main thesis of this chapter: disinformation now operates in a coordinated configuration of text, image, sound, and video, so analysis must also be based on a multimodal, ecosystem-level methodology.

Subchapter synthesis (short): deepfakes and synthetic media are not just “fake content,” but a crisis of evidence status; the concept of credibility hijacking links this change to the WGC framework.

Solution attempt: the ecosystem model of “weapon-grade communication” (WGC) and the WGC matrix (introduction of the new method)

Based on the diagnosis in the previous chapters, the problem is not “AI knowledge” but the socio-technical configuration of AI’s

use as a weapon: the combination of actors, goals, channels, and effects, which is simultaneously a phenomenon of communication theory, criminology, and military science. Therefore, the WGC ecosystem model (Horváth, 2023c) is not a new taxonomy, but an analytical order: it breaks down the same phenomenon at the ecosystem–operational–tactical levels and organizes the interpretation specifically around the security policy question “what is at stake?”

The WGC matrix:

the intersection of intent, channel, and cognitive effect

The purpose of the WGC matrix is to view AI-supported content not only as content (what does it claim?), but also as an operational product (why, how, with what consequences?). The three axes of the matrix are therefore: (i) intent (what is the intervention aimed at: destabilization, undermining legitimacy, intimidation, making money, etc.), (ii) channel configuration (open platforms, closed networks, advertising ecosystem, bot/fake profile infrastructure), and (iii) cognitive impact (what psychological and social consequences does it push towards: erosion of trust, paralysis of action, panic, cynicism). With this logic, agenda-setting or the “firehose of falsehood” is not a separate theoretical block, but a description of the intention–channel–effect relationship (see: McCombs–Shaw, 1972; Paul–Matthews, 2016).

While the S–F–C (Situation–Forecast–Control) grid describes the temporal and

management dynamics of cognitive warfare, the WGC matrix introduced in this study performs the structural operationalization of the starting point of this process, the Situation. The WGC matrix thus functions as a diagnostic "magnifying glass": it breaks down a static snapshot of the security environment (S) into dynamic operational elements (intent, channel configuration, cognitive effect), thereby laying the foundation for subsequent forecasting (Forecast) and intervention (Control).

According to the short protocol of methodological operationalization, coding is performed on the three axes of the WGC matrix (intent; channel configuration; cognitive effect). This chapter operationalizes two key outputs with indicator groups: (i) credibility hijacking, (ii) cognitive overload. The rating is done on a three-point scale (low/medium/high), and each value is linked to a documented source chain (platform report/cybersecurity report/OSINT/fact-check/official announcement). We use negative testing: if the intent or channel coordination cannot be substantiated by the source chain, the rating is downgraded or the claim is not upheld.

Operationalization of credibility hijacking and cognitive overload

The matrix works when the "effect" is not an impression but a consistently marked group of indicators. This study highlights two WGC outputs that are particularly common in

multimodal environments: credibility hijacking and cognitive overload.

Authenticity deviation refers to solutions that short-circuit the recipient's "sense of evidence" (e.g., iconic images, "evidential" videos, authoritative voices, institutional branding), while source reconstruction deteriorates; This fits in with the source monitoring theory, according to which source traces fade faster than content memory (Johnson, Hashtroudi, Lindsay, 1993).

By cognitive overload, I mean a state of quantitative scaling (variants, languages, repetition, channel synchronization) in which the recipient encounters not a statement but a persistent noise environment; Therefore, labeling does not occur on a "true/false" axis, but rather by recording the pattern of propagation, repetition, variant formation, and channel crossover.

As a practical designation (to save characters), three-level ratings can be given: low/medium/high for a) intensity of credibility hijacking, b) overload pressure, c) channel coordination; in each case, a source chain is linked to the categories.

Reliability and validity of the method based on case studies

The reliability of the method is ensured by the fact that the case studies are based on the same analytical framework: the same questions must be answered (actor–intention–channel–effect–reaction) and coded using the same coding rules (Horváth, 2023c). Validity here is not a laboratory result, but a chain of

evidence: every statement is backed by a traceable chain of sources (platform report / cybersecurity report / OSINT / fact-check / official announcement), and the matrix also forces negative testing: if the intention cannot be verified or the channel coordination cannot be substantiated, the rating is downgraded. This also implies the practical meaning of falsifiability: the goal is not to "prove the WGC," but to find out where it does not hold.

Validation: case study-based application – selection logic and analytical framework

Validation is based on case studies: selection is not aimed at representativeness, but at methodological coverage. The logic behind this is: (i) multimodal spectrum (text–image–sound–video). The order of the main categories follows the increase in technological complexity and psychological impact: we move from one-dimensional signals requiring cognitive processing (text) to multimodal forms aimed at visceral perception and unconditional trust (video). (ii) different sets of actors (state influence and criminal exploitation), (iii) different stakes (electoral integrity, critical infrastructure, public trust), (iv) documentation (public reports and verifiable source chain). The next chapter is therefore not a simple "list of cases," but a test of the matrix. The order of the case studies within each category reflects the evolution of the threat: we move from initial technological demonstrations or individual-level "noise" towards system-level operations that pose a strategic risk.

Text generated by artificial intelligence

Case 1: The evolution of the Dragonbridge network from revenge porn to high politics: how Chinese spam became a strategic weapon

The Dragonbridge/Spamouflage prototype is an example of how a low-quality, "spam-like" text stream can become a persistent, ecosystem-level weapon of communication: it does not seek to win a single argument, but rather shapes the background noise of the digital public sphere. The network was described by Graphika as Spamouflage Dragon (Graphika, 2020) and later documented by the Google Threat Analysis Group as DRAGONBRIDGE, with massively removed channels and multilingual content (Google TAG, 2022; Google TAG, 2024); state interests were also confirmed by cybersecurity reports (Mandiant, 2022), and transnational pressure patterns also emerged (Rapid Response Mechanism Canada, 2023–2025). The role of AI here is scaling: translation, variation, rapid reproduction of templates, and "circumventing" moderation filters with many small narrative variations (Google TAG, 2024). In the WGC section, the intention is long-term information environment shaping; the channel is a multiplatform, multilingual, loosely connected network of accounts; the cognitive effect is uncertainty/erosion of trust through cognitive overload and thematization (see: McCombs–Shaw, 1972).

Distribution: YouTube/Facebook/X/TikTok + blogs/forums; massive, varied posts and comments, even with low organic reach, as "constant background noise."

Coding: Intent = strat. thematization/erosion; Channel = multi-platform–multi-language–fake profile network; Cognitive effect = overload/apathy/erosion of trust; Coordination = medium–high.

Case 2: “Hi Mom, I’m in trouble!” – a grammatically flawless attempt at cheating, or the role of generative AI in digital bank robbery

The AI-based phishing prototype is a case where weapon-grade communication is not spectacular at the campaign level, but rather “scatters” and destroys digital trust at the micro level. According to several recent studies, the persuasiveness of spear phishing messages written by generative models can match or even exceed that of human experts, while the cost and time of production are radically reduced (arXiv, 2024; Electronics, 2024; MDPI, 2025). AI plays a threefold role here: error-free text writing tailored to the brand voice; variant generation (against filters); personalization based on OSINT/previous leaks. In WGC terms, the intention is to gain access/money in the short term and to erode crisis communication and institutional credibility in the long term; the channel is “everyday” infrastructure (e-mail/SMS/messaging/private messages); the cognitive effect is based on urgency, the exploitation of automatisms based on authority mimicry and trust-building, followed by the establishment of lasting suspicion. In terms of national security, the method can slip into the realm of hack-and-leak if it targets institutions/elections (Canadian Centre for Cyber Security, 2024).

Distribution: email + SMS + messenger + private social media messages; brand imitation, urgent “identification/update” narratives, redirection to clone sites (OTP Bank, 2024).

Coding: Intent = access/money → erosion of trust; Channel = multi-channel direct message; Cognitive effect = urgency/authority/automatic trust → general suspicion; Coordination = low–medium (depending on campaign type).

Case 3: The anatomy of operation Doppelgänger: when europe’s most read newspapers begin to spread russian propaganda

Doppelgänger is a type of textual WGC where the weapon is not quantity but the appearance of credibility: the user reads a pro-Russian narrative in a “familiar media voice.” The campaign was named and described by EU DisinfoLab after a network of cloned sites imitating European media outlets was identified in 2022 (EU DisinfoLab, 2022–). The sustainability of the operation is supported by investigative and government analyses: fake articles are published on a daily basis across multiple channels; USCYBERCOM has also described the structure as complex web deception (Correctiv, 2024; USCYBERCOM, 2024). The role of AI here is “style and language alignment”: large quantities of quickly produced texts that conform to journalistic style and only shift the narrative at key points (sanctions, support, war fatigue). In WGC terms, the intention is to undermine European cohesion and the legitimacy of NATO/EU policies; the

channel is domain and brand cloning + targeted social media dissemination; the cognitive effect is to disrupt source monitoring and reinforce the "media is not reliable" narrative.

Distribution: cloned news sites + X/Facebook accounts + closed messaging groups; link sharing with search engine and sharing-optimized texts.

Coding: Intent = legitimacy destruction/cohesion weakening; Channel = domain/brand cloning + targeted distribution; Cognitive effect = source confusion/trust erosion/war fatigue; Coordination = high.

Subchapter synthesis (short): The common lesson of text-based case studies is that generative AI weaponizes the nature of "text as infrastructure": it not only produces content, but also creates a scalable, variable, and multilingual narrative presence that places a lasting strain on the public's attention and monitoring capacity.

Dragonbridge/Spamouflage exemplifies this model in terms of quantitative superiority and multilingual variant production ("low quality, high volume" background noise), while Doppelgänger is its counterpart: here, the decisive resource is not quantity, but authenticity diversion (brand/domain cloning, press style imitation), i.e., confusing source identification. Generative phishing of the "Hi Mom..." type completes the picture from the perspective of micro-level persuasion: with the disappearance of linguistic errors, the former lay filters (poor spelling, foreign expressions) lose their

validity, and the erosion of trust seeps into everyday transactional communication. Together, the three cases illustrate the chain whereby textual WGC is capable of simultaneously shaping the agenda, falsifying sources, and triggering direct action (clicking/referral/abstention), while the long-term "gain" is the erosion of trust through cognitive overload.

Image generated by artificial intelligence

Case 4: Puffer-Jacket Balenciaga Pope:: how he convinced millions with a single AI-mem that Pope Francis had changed his style

The case of the "Balenciaga Pope" prototype is one where the illusion of photorealism becomes an independent WGC component: the image does not state a "fact," yet it massively activates the reflex of visual evidence. The visual was launched in March 2023 (Reddit) and then transferred to major platforms with a loss of context; many interpreted it as a press photo, and recognition of its generated origin typically only occurred after it went viral (CBS News, 2023; The Guardian, 2023; Reuters Fact Check, 2023). The role of AI here is twofold: to produce a photo effect with a low entry threshold and to blur the line between meme and news photo, which in the long run weakens the norm of the photo as evidence (Horváth, 2023).

Summary (WGC brief): Evidence: viral spread + fact checks. Role of AI: photo effect from prompt, acceleration of context loss. Intent: attention/thematization (without harm, but

with a norm-destroying effect). Channel: platform feeds. Cognitive effect: visual evidence heuristics, delayed suspicion. Distribution: Reddit → X/Facebook/Meta. Coordination: low.

Case 5: A few seconds of stock market chaos that wiped out hundreds of billions of dollars with a fake photo

The “explosion photo” near the Pentagon is no longer a pop culture anomaly, but a condensed visual shock that activates market and security reflexes: a single image caused a brief real market turmoil before the refutation stabilized (The Guardian, 2023; Al Jazeera, 2023). The visual most likely originated from a generative image system; the news photo scheme worked, while typical artifacts appeared in the details (UC Berkeley School of Information, 2023). The function of AI here is low-cost panic-mongering: a quickly produced, believable composition that triggers heuristics (danger, urgency, national security) even without a narrative (Paul and Matthews, 2016). At the WGC level, the intent is to distort risk perception in the short term; the channel is platform credibility markers and accounts posing as “news sources”; the cognitive effect is to increase the likelihood of false alarms and overreactions.

Summary (WGC brief): Evidence: official refutation + press analyses + artifact notes. AI role: generating news photo patterns to create “shock.” Intent: distortion of risk perception, activation of panic reflex. Channel: credibility markers and viral network. Cognitive effect:

false alarm, overreaction, short-term uncertainty. Dissemination: X → media coverage → refutation. Coordination: moderate.

Case 6: Millions of shares for the “perfect” war photo: when reality is too bloody, we draw it pretty

The “All Eyes on Rafah” case is the third stage of generated images: it is not classic disinformation, yet it is WGC -relevant because it scales the toolkit of moral framing and reputational pressure. In May 2024, a declared AI-generated visual spread widely and briefly became a “mandatory” reference point in debates about the humanitarian dimension of the conflict (Al Jazeera, 2024). The power of the image lies not in its replication of reality, but in its “symbolic condensation” of real suffering: it is shareable, emotionally charged, and, through its story format, elevates the visual space of private profiles to the conflict agenda. Debates about the blurring of the line between photo and illustration and the issue of credibility costs have intensified in retrospect (ABC, 2024). At the WGC level, the intention is to shift the moral agenda; the channel is the platform ecosystem and the remixable template; the cognitive effect is normative coercion and source forgetting, which can be quickly followed by counter-images and counter-campaigns (Newsweek, 2024).

Summary (WGC brief): Evidence: platform spread data + press debate. AI role: symbolic condensation, remixability, rapid global scaling. Intent: moral framing, mobilization, reputational pressure. Channel: Meta/X/

TikTok, especially Instagram stories. Cognitive effect: normative coercion, source amnesia, credibility cost risk. Distribution: mass story sharing, cross-platform remixing. Coordination: low–medium.

Subheading synthesis (short): The common denominator of image-based case studies is that generative AI, through photorealism, begins to challenge the status of "visual evidence": the viewer believes what they see before checking the source. The puffy-jacketed Balenciaga pope is the threshold case: as a pop-cultural, "harmless" viral, he normalizes the fact that a believable photo is no longer a guarantee, thus paving the way for later, more targeted operations. The fake photo of the Pentagon explosion already demonstrates the systemic consequences: a single visual impulse can activate market and security reflexes for a short time, thus radically reducing the cost of image-based panic. All Eyes on Rafah shows the third dimension: the generated image does not necessarily function as a "lie," but with moral framing and condensation into a memetic symbol, it can shape reputational pressure and the agenda, while forgetting the source and blurring the line between photo and illustration can generate credibility costs. Together, the three cases indicate that visual WGC is not just about falsification: the decisive factor is that visual content spreads faster than correction and, in the long run, rewrites the routines of controlling the collective image of reality.

Artificial intelligence-generated voice

Case 7: The €220,000 phone call when the boss's voice orders a bank robbery

The case of CEO voice cloning demonstrates the "precision" use of WGC against micro-level, high-value targets: a single call can activate organizational obedience patterns. In a case that made headlines, fraudsters imitated the voice of a senior executive at an energy company and instructed the target to make an urgent transfer; the voice (accent, intonation, familiar tone) was enough to carry out the transaction. Artificial intelligence here "arms" classic social engineering: the source of communication (the boss) becomes falsifiable, and thus deception occurs through the most confidential channel. Its military relevance arises where the corporate decision-making hub is linked to critical infrastructure: a wrong decision can lead not only to financial losses but also to operational disruptions and supply chain distortions, and thus converge with state/quasi-state interests (Horváth, 2023; Horváth, 2024).

Summary (WGC brief): Evidence: incident description + cybersecurity interpretation. AI role: voice model from a short sample, reproduction of a "command" voice. Intent: money/access → distortion of critical decisions. Channel: targeted phone call (outside moderation visibility). Cognitive effect: authority and urgency reflex, activation of organizational obedience. Distribution: not mass; individual, targeted calls. Coordination: low–medium.

Case 8: Weaponizing telephone harassment: "Don't go vote!" – said the Biden clone voice

The 2024 Biden robocall case shows how AI-generated voice crosses from the experimental space into the WGC domain, directly threatening election integrity. Ahead of the New Hampshire primary, automated calls imitating the voice of Joseph R. Biden urged voters to "not go vote" but to "save their vote for November." The voice (tone, rhythm, intonation) acted as an anchor of authority: the recipient tends to identify the recognizable voice with authenticity, even if the content itself is suspicious. The technological innovation is not the robocall as a channel, but the scalability of voice cloning: short sample → reproducible "presidential voice" → mass, variable messaging with minimal human input. In terms of communication theory, this is the voice-based equivalent of action-oriented, repeated messaging (firehose logic); in criminological terms, it is an attempt to manipulate voter behavior; and from a military science perspective, it can be linked to hybrid goals of destabilizing democratic institutions (see: McCombs–Shaw, 1972; Paul and Matthews, 2016; Horváth, 2022; Horváth, 2024).

Summary (WGC brief): Evidence: official investigation + press/fact-check confirmation. AI role: voice cloning and scalable robocall variants. Intent: to reduce participation, undermine the process. Channel: telephone network + automated calling system. Cognitive effect: authority reflex, false authenticity, timing pressure. Dissemination:

robocall → online press + social media reaction/refutation. Coordination: moderate.

Case 9: 48 Hours before the polls close: the audio recording that poisoned Slovakia

The deepfake audio recording before the 2023 Slovak parliamentary elections is an example of targeted, regional application: even in a small country, it can interfere with the electoral process if the timing follows the logic of an "October surprise." Close to the campaign silence period, an audio recording was released in which Michal Šimečka and Monika Tódová allegedly discuss "manipulating" the election; its rapid spread was later interpreted by fact-checking and expert analyses as AI-generated manipulation. The technical role here is no longer merely the reproduction of sound: authenticity is reinforced by the imitation of discursive style (phrasing, tempo, "local" speech patterns), i.e., the model produces a "character" who credibly enters into a politically charged dialogue. The core of communication theory is the undermining of source credibility (what did I actually hear?), which quickly generates confusion and distrust; criminologically, this is targeted intervention in the electoral process; militarily, it is the "contamination" of the Central and Eastern European information space (McCombs–Shaw, 1972; Paul and Matthews, 2016; UNESCO, no year).

Summary (WGC brief): Evidence: rapid viral spread + fact-checking + expert evaluations. AI role: voice cloning + discursive style imitation. Intent: undermining trust, weakening legitimacy, shock before the campaign

silence period. Channel: social media spread + forwarding chains. Cognitive effect: source uncertainty, reputational damage, rapid polarization. Distribution: platforms + messengers, forwarding in closed groups. Coordination: medium–high.

Subchapter summary (short): The common denominator of the three voice cases is that the attack does not require a new platform, but rather weaponizes everyday voice channels: robocalls, social media forwarding, targeted phone calls. The medium of voice is critical because the heuristic of “recognizable voice = authenticity” kicks in faster than source verification, especially under time pressure. As a result, refutation is often structurally delayed: some of the damage/distortion occurs before the correction.

Artificial intelligence-generated video

Case 10: This is not Morgan Freeman: demonstration deepfake and the proof video crisis

The “This is not Morgan Freeman” demonstration deepfake video is the “zero milestone” of moving image manipulation: it is not an operational campaign, but a technological demonstration that simultaneously teaches the audience that “AI can already do this” and normalizes the idea that faces and voices can be reproduced at any time. The video declares from the outset that we are watching a deepfake (Avid Open Access, n.d.; HP, 2024), yet its effect is WGC-like: it shatters the layperson's sense of

security that video and audio are “proof.” Subsequent empirical results support this phenomenon: the debate surrounding deepfakes is not only technical, but also concerns a crisis of “epistemic trust” (Twomey, 2023), and we perform poorly in recognizing them even with instructions (Somoray, Miller, 2023).

In terms of WGC, this case is not about direct influence, but about changing the environment: according to the logic of “shaping the information environment,” social perception thresholds and credibility criteria are rewritten, making subsequent political/commercial/criminal applications more effective (McCombs, Shaw, 1972; Paul, Matthews, 2016; Horváth, 2023).

Summary (WGC brief): Evidence: demonstration video + professional/educational references + empirical recognition results. AI role: photorealistic face + voice cloning + lip-syncing. Intent: normalization/“showcasing capabilities” → preparing for the erosion of credibility standards. Channel: public platforms (YouTube, professional portals, embeds). Cognitive effect: weakening of the status of the video as evidence, epistemic uncertainty. Distribution: mostly organic. Coordination: low.

Case 11: Karikó Katalin on the crypto exchange: deepfake advertisements causing billions in damages

The second pattern of deepfake videos is the world of investment and crypto fraud: fake offers are “legitimized” with the faces and voices of well-known public figures. According to cybersecurity and regulatory

reports, these campaigns are increasingly operating with AI-based videos, linked to organized boiler room infrastructure (Unit 42, 2024). Official warnings and reports also confirm this phenomenon (e.g., Government of Western Australia, 2025; ACCC Scamwatch, 2024), pointing out that video is no longer just an illustration, but the main tool for gaining trust.

At the technical-operational level, this is a multimodal pipeline: LLMs write the text variants, voice cloning reproduces the "expert" voice, and the deepfake engine puts together the facial expressions and synchronization. The "firehose" also works at the format level here: the same lie returns as a video, post, story, or advertisement (Paul, Matthews, 2016; Horváth, 2023). Criminologically, the pattern is linked to organized crime, money laundering chains, and call centers; at the systemic level, it undermines trust in financial communication. Its military relevance lies in the fact that it increases the vulnerability of critical financial infrastructures to information attacks and can generate longer-term instability.

Summary (WGC brief): Evidence: cybersecurity/authority reports + identified network patterns. AI role: cloning of celebrity/authority faces and voices + text variant generation + multimodal coordination. Intent: making money and gaining access → erosion of trust in the financial sector. Channel: video ads and short videos (Meta, TikTok, YouTube, messaging apps). Cognitive impact: authority heuristic, cognitive overload, "everything is suspicious"

trust defense. Distribution: repetitive campaigns reinforced by recommendation systems. Coordination: high.

Case 12: Even poor-quality lies can kill: anatomy of Zelensky's fake "Lay down your arms!" video

The deepfake attributed to Zelensky (March 2022) made the military science stakes of moving image manipulation explicit: this is no longer a game of reputation, but an attempt to distort the chain of command, morale, and situation assessment (see: Euronews, 2022; Smalley, 2022). Although technically immature, it went down in history as the first deepfake video attributed to a head of state to appear in an active war zone. The key to the chronology is resilience: on the Ukrainian side, pre-bunking warnings, rapid debunking, and an authentic refutation video (The Guardian, 2022), which reduced the immediate impact, while at the same time the example immediately entered international professional and legal discourse (Kuźnicka-Błaszowska, 2025).

In the WGC framework, the video combines the TEXT and IMAGE case types: the narrative of surrender + iconic visual frame (presidential setting, familiar background) + the illusion of moving image authenticity together target the most sensitive audience. The relevant lesson is that early, poor quality is not harmless: technology improves quickly, but the social immune system is slower (Desjardins, 2021).

Summary (WGC brief): Evidence: incident reports + removal/refutation + literature/legal processing. AI role: face and facial

expression manipulation + lip-syncing + imitation of the visual frame of a "presidential speech." Intent: surrender/resistance narrative → undermining morality and command structure. Channel: hacked news portals + social platforms (Facebook/X) + possible stream slippage. Cognitive effect: misjudgment of the situation, panic/breakdown, erosion of trust. Distribution: multi-channel, with rapid removal. Coordination: medium–high.

Subheading synthesis (short): The three video cases follow the same pattern: (1) demonstration deepfakes as a "basic motif" of normalization (crisis of evidence videos), (2) deepfake advertisements used in organized crime as industrial scaling (multimodal firehose), (3) war deepfakes as direct operational relevance (morale/command/credibility). Their common WGC point is that moving images simultaneously convey the illusion of visual evidence and audio authenticity, so refutation is structurally delayed: the effect often runs its course before debunking, even if the manipulation is easily recognizable in hindsight.

Conclusion: new ways of defense

The essence of multimodal influence accelerated by generative AI is not merely "more false statements," but a structural shift in the ecosystem of attention and credibility: the agenda (what we talk about), dissemination (how it reaches us), and source perception (who we believe) all become manipulable at

once (McCombs, Shaw, 1972; Paul, Matthews, 2016; Johnson et al., 1993). In this framework, the new way to defend ourselves is not a single "better detector," but an analytical and decision-support approach that measures risk at the intersection of intent, channel, and cognitive effect, thus capturing it not in binary terms (true/false) but in terms of operational relevance.

What does the matrix offer that detection-centric thinking does not?

The detection-focused approach typically responds to the "authenticity" of content, while the WGC matrix focuses on the function of the attack: what it aims to achieve, through which channel ecosystem, and what cognitive vulnerability it targets (Paul, Matthews, 2016; Wardle, Derakhshan, 2017). The framework therefore works even when the attack is not a "classic lie" but rather context and identity falsification or "evidential media" authenticity diversion (Chesney, Citron, 2019; Johnson et al., 1993). The added value is practical: it makes cases with different modalities (text–image–sound–video) comparable and helps to plan defenses not at a single point (content control) but across the entire chain (distribution logic, attention cycle, source labeling, institutional response time) (Bakshy et al., 2015; Fazio et al., 2015).

Critical infrastructure, election integrity, public trust: implications for application

In the case of critical infrastructure, the stakes are often not mass persuasion, but confidence-eroding disruption: even a small

amount of "evidential" content can be enough to cause customer panic, service and call center overload, and reputational damage. The advantage of the matrix is that it operationalizes the damage mechanism (cognitive overload, source amnesia, illusory truth effect) from early signs, rather than treating it as a retrospective narrative explanation (Fazio et al., 2015; Johnson et al., 1993).

In election integrity, "last mile" vulnerability—especially voice and video manipulation with a short reaction window—is a risk where legal and platform-level frameworks (e.g., DSA) are necessary but not sufficient on their own: the decisive factors are operational timing and psychological impact management, i.e., the ability to quickly pre-bunk/debunk and manage channel coordination (EU, 2022; UNESCO, 2024).

At the level of public trust, the challenge is to ensure that the public does not merely doubt individual statements, but also the minimum conditions of shared reality. The matrix therefore also "translates" for strategic decision-making: it classifies the content event as a security risk and identifies where the greatest systemic damage is likely to occur (Horváth, 2022; Horváth, 2024, 2025).

Limitations and further research directions

The proposed matrix is not a detector test or an automated truth-checking system: it is based on documented sources, a reconstructable source chain, and transparent coding, and therefore depends on data availability and

analyst consistency (Wardle, Derakhshan, 2017). A further limitation is that direct measurement of cognitive effects (e.g., source amnesia, illusory truth) is often only possible with proxy indicators, which requires strict validation discipline (Johnson et al., 1993; Fazio et al., 2015).

The next steps should be: (i) developing inter-coder reliability procedures and a labeling standard supported by a list of examples, (ii) comparing the matrix with platform regulation compliance logics and case types (EU, 2022), and (iii) expanding the targeted, comparable case study corpus in the CI-choice-public trust triangle so that the model can maintain falsifiable statements in different operational environments (see: UNESCO, 2024; Europol, 2022).

References

- ABC News (2024). What is 'All Eyes on Rafah'? *ABC News* (Australia).
Downloaded: 2025.11.26. Web:
<https://www.abc.net.au>
- ACCC Scamwatch (2024). *Investment scams involving deepfakes*. ACCC Australia.
Downloaded: 2025.11.26. Web:
<https://www.scamwatch.gov.au>
- Al Jazeera (2023). Fake Pentagon explosion photo goes viral, markets dip. *Al Jazeera*.
Downloaded: 2025.11.26. Web:
<https://www.aljazeera.com>
- Al Jazeera (2024). 'All Eyes on Rafah': The AI image shared by millions. *Al Jazeera*.
Downloaded: 2025.11.26. Web:
<https://www.aljazeera.com>

- Andrew, C. és Mitrokhin, V. (2000). *The Sword and the Shield: The Mitrokhin Archive and the Secret History of the KGB*. Basic Books, New York.
- arXiv (2024). *Generative AI and Phishing / Spear-phishing studies*. arXiv preprint. Downloaded: 2025.11.26. Web: <https://arxiv.org>
- Avid Open Access (n. d.). *This is not Morgan Freeman - Deepfake demonstration*. Avid Open Access. Downloaded: 2025.11.26. Web: <https://avidopenaccess.org>
- Bakshy, E., Messing, S. és Adamic, L. A. (2015). Exposure to ideologically diverse news and opinion on Facebook. *Science*, 348(6239), 1130–1132. DOI: <https://www.doi.org/10.1126/science.a1160>
- Barron, J. (1983). *KGB Today: The Hidden Hand*. Reader's Digest Press, New York.
- Bernays, E. L. (1928). *Propaganda*. Horace Liveright, New York.
- Canadian Centre for Cyber Security (2024). *Cyber threats to democratic processes*. Government of Canada. Downloaded: 2025.11.26. Web: <https://cyber.gc.ca>
- CBS News (2023). Pope Francis 'puffer jacket' fake photo goes viral. CBS News. Downloaded: 2025.11.26. Web: <https://www.cbsnews.com>
- Chesney, R. és Citron, D. K. (2019). Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, 107(6), 1753–1819. DOI: <https://www.doi.org/10.15779/Z38RV0D15J>
- Correctiv (2024). *Doppelgänger investigation*. Correctiv.org. Downloaded: 2025.11.26. Web: <https://correctiv.org>
- Couldry, N. és Hepp, A. (2017). *The Mediated Construction of Reality*. Polity Press, Cambridge.
- Desjardins, J. (2021). The rapid evolution of deepfake technology. *Visual Capitalist*. Downloaded: 2025.11.26. Web: <https://www.visualcapitalist.com>
- Electronics (2024). AI in Phishing detection and generation. *Electronics*, 13(5). DOI: <https://www.doi.org/10.3390/electronics13050000>
- EU DisinfoLab (2022). Doppelgänger - Media Cloning Campaign. *EU DisinfoLab Report*. Downloaded: 2025.11.26. Web: <https://www.disinfo.eu>
- Euronews (2022). Deepfake video of Zelensky calling on Ukrainians to surrender. *Euronews*. Downloaded: 2025.11.26. Web: <https://www.euronews.com>
- European Union (2022). Digital Services Act (DSA) – Regulation (EU) 2022/2065. *EUR-Lex*. Downloaded: 2025.11.26. Web: <https://eur-lex.europa.eu>
- Europol (2022). Facing the Challenge of Deepfakes in Cybercrime and Disinformation. Europol Innovation Lab, *The Hague*. Downloaded: 2025.11.26. Web: <https://www.europol.europa.eu>
- Fazio, L. K., Brashier, N. M., Payne, B. K. és Marsh, E. J. (2015). Knowledge Does Not Protect Against the Illusory Truth Effect. *Journal of Experimental Psychology: General*, 144(5), 993–1002. DOI:

- <https://www.doi.org/10.1037/xge0000098>
- Gillespie, T. (2018). *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions That Shape Social Media*. Yale University Press, New Haven.
- Google Threat Analysis Group (2022). *DRAGONBRIDGE activity. Google TAG Report*. Downloaded: 2025.11.26. Web: <https://blog.google/threat-analysis-group>
- Google Threat Analysis Group (2024). *DRAGONBRIDGE updates. Google TAG Report*. Downloaded: 2025.11.26. Web: <https://blog.google/threat-analysis-group>
- Government of Western Australia (2025). *ScamNet alerts on deepfake investment scams*. WA Government. Downloaded: 2025.11.26. Web: <https://www.scamnet.wa.gov.au>
- Graphika (2020). Spamouflage Dragon: Spamouflage network. *Graphika Report*. Downloaded: 2025.11.26. Web: <https://graphika.com>
- Hariman, R. és Lucaites, J. L. (2007). *No Caption Needed: Iconic Photographs, Public Culture, and Liberal Democracy*. University of Chicago Press, Chicago.
- Hoffman, F. G. (2007). *Conflict in the 21st Century: The Rise of Hybrid Wars*. Potomac Institute for Policy Studies, Arlington.
- Horváth, D. (2022). Pandemic and Infodemic – Which One Is More Dangerous? *Acta Universitatis Sapientiae, Communicatio*, 9(1), 125–139. DOI: <https://www.doi.org/10.2478/auscom-2022-0009>
- Horváth, D. (2023a). Az álhír fogalma és helye a hibrid hadviselésben. *Nemzet és Biztonság*, 2023/1, 45–60.
- Horváth, D. (2023b). Az álhír és a hibrid hadviselés kapcsolata. In *A hadtudomány és a 21. század*. Ludovika Egyetemi Kiadó, Budapest. pp. 88–102.
- Horváth, D. (2023c). A fegyvernek minősülő kommunikáció (FMK) ökoszisztéma-modellje és FMK-mátrix. *Hadtudományi Szemle*, 16(3), 21–45.
- Horváth, D. (2024). Az álhír, vagyis a fegyvernek minősülő kommunikációs taktika mint aktuális hadtudományi kihívás. In Gazdag, F., Padányi, J. és Tóth, P. (eds.). *A hadtudomány aktuális kérdései* 2023. Ludovika Egyetemi Kiadó, Budapest. pp. 151–168.
- Horváth, D. (2025). A COVID–19-válság szerepe a fegyvernek minősülő kommunikáció új dimenzióinak kialakulásában. In *Lélektan és hadviselés – interdiszciplináris folyóirat*, 2025/2. szám (megjelenés alatt).
- HP (2024). Deepfake risks and demonstrations. *HP Wolf Security*. Downloaded: 2025.11.26. Web: <https://threatresearch.ext.hp.com>
- Johnson, M. K., Hashtroudi, S. és Lindsay, D. S. (1993). Source monitoring. *Psychological Bulletin*, 114(1), 3–28. DOI: <https://www.doi.org/10.1037/0033-2909.114.1.3>
- Jowett, G. S. és O'Donnell, V. (2019). *Propaganda & Persuasion*. SAGE Publications, Thousand Oaks.

- Kalugin, O. (1994). *The First Directorate: My 32 Years in Intelligence and Espionage Against the West*. St. Martin's Press, New York.
- Kirchenbauer, J. és tsai (2023). *A Watermark for Large Language Models*. arXiv preprint. Downloaded: 2025.11.26. Web: <https://arxiv.org/abs/2301.10226>
- Kramer, F. D. és Speranza, L. D. (2017). *Meeting the Russian Hybrid Challenge*. Atlantic Council, Washington D.C.
- Kuźnicka-Błaszowska, D. (2025). Legal implications of wartime deepfakes. *International Journal of Law and IT* (Kézirat).
- Lasswell, H. D. (1927). *Propaganda Technique in the World War*. K. Paul, Trench, Trubner & Co., London.
- Lasswell, H. D. (1948). The Structure and Function of Communication in Society. In Bryson, L. (Ed.). *The Communication of Ideas*. Harper and Row, New York. pp. 37–51.
- Mandiant (2022). DRAGONBRIDGE/Spamouflage activity. *Mandiant Threat Intelligence*. Downloaded: 2025.11.26. Web: <https://www.mandiant.com>
- McCombs, M. E. és Shaw, D. L. (1972). The Agenda-Setting Function of Mass Media. *Public Opinion Quarterly*, 36(2), 176–187. DOI: <https://www.doi.org/10.1086/267990>
- MDPI (2025). Emerging threats in AI phishing. *Future Internet*, 17(1). DOI: <https://www.doi.org/10.3390/fi17010000>
- Mitchell, W. J. T. (1994). *Picture Theory: Essays on Verbal and Visual Representation*. University of Chicago Press, Chicago.
- Morozov, E. (2013). *To Save Everything, Click Here: The Folly of Technological Solutionism*. PublicAffairs, New York.
- NATO (2016). *Hybrid Warfare / Hybrid Threats*. NATO Archives. Downloaded: 2025.11.26. Web: <https://www.nato.int>
- Newsweek (2024). 'All Eyes on Rafah' AI image sparks controversy. *Newsweek*. Downloaded: 2025.11.26. Web: <https://www.newsweek.com>
- OTP Bank (2024). *Tájékoztatás adatahalász kísérletekről*. OTP Bank hivatalos közlemény. Downloaded: 2025.11.26. Web: <https://www.otpbank.hu>
- Paul, C. és Matthews, M. (2016). *The Russian "Firehose of Falsehood" Propaganda Model: Why It Might Work and Options to Counter It*. RAND Corporation, Santa Monica. DOI: <https://www.doi.org/10.7249/PE198>
- Rapid Response Mechanism Canada (2023–2025). *RRM Reports on Foreign Interference*. Government of Canada. Downloaded: 2025.11.26. Web: <https://www.canada.ca>
- Reuters Fact Check (2023). Fact Check: Image of Pope Francis in puffer jacket is AI generated. *Reuters*. Downloaded: 2025.11.26. Web: <https://www.reuters.com>
- Rid, T. (2020). *Active Measures: The Secret History of Disinformation and Political Warfare*. Farrar, Straus and Giroux, New York.
- Shifman, L. (2014). *Memes in Digital Culture*. MIT Press, Cambridge.
- Smalley, I. (2022). *The first weaponized deepfake in war?* Tech Policy Press. Downloaded: 2025.11.26. Web: <https://techpolicy.press>

- Somoray, K. és Miller, D. J. (2023). Detection of deepfake videos. *Computers in Human Behavior*, 145. DOI: <https://www.doi.org/10.1016/j.chb.2023.107765>
- The Guardian (2022). Deepfake of Zelensky surrender removed from social media. *The Guardian*. Downloaded: 2025.11.26. Web: <https://www.theguardian.com>
- The Guardian (2023). Fake photo of Pentagon explosion causes brief market dip. *The Guardian*. Downloaded: 2025.11.26. Web: <https://www.theguardian.com>
- The Guardian (2023). Pope Francis in a puffer jacket: AI-generated image fools the internet. *The Guardian*. Downloaded: 2025.11.26. Web: <https://www.theguardian.com>
- Twomey, J. (2023). The Epistemic Threat of Deepfakes. *Synthese*, 201(3). DOI: <https://www.doi.org/10.1007/s11229-023-04000-x>
- UC Berkeley School of Information (2023). *Analysis of the fake Pentagon explosion image artifacts*. UC Berkeley. Downloaded: 2025.11.26. Web: <https://www.ischool.berkeley.edu>
- UK Ministry of Defence (2013). *Joint Doctrine Note 2/13: Information Operations*. Ministry of Defence, Shrivenham. Downloaded: 2025.11.26. Web: https://assets.publishing.service.gov.uk/media/5a7bf3f3ed915d74e33f2cbb/20130828-JDN_2_13_Information_Operations.pdf
- UNESCO (2024). *Elections and the Threat of AI-generated Disinformation: Case Study – Slovak Parliamentary Elections 2023*. UNESCO. Downloaded: 2025.11.26. Web: <https://www.unesco.org>
- Unit 42 (2024). Deepfake scams and boiler room operations. *Palo Alto Networks / Unit 42*. Downloaded: 2025.11.26. Web: <https://unit42.paloaltonetworks.com>
- USCYBERCOM (2024). Complex Web Deception / Doppelgänger. *US Cyber Command*. Downloaded: 2025.11.26. Web: <https://www.cybercom.mil>
- Wardle, C. és Derakhshan, H. (2017). *Information Disorder: Toward an interdisciplinary framework for research and policymaking*. Council of Europe, Strasbourg. Downloaded: 2025.11.26. Web: <https://rm.coe.int/information-disorder-toward-an-interdisciplinary-framework-for-research/168076277c>
- Winner, L. (1980). Do Artifacts Have Politics? *Daedalus*, 109(1), 121–136.

KREATÍV TÖRTÉNETMESÉLÉS OZOBOTOK SEGÍTSÉGÉVEL

Szerző:

Szabó Tibor (Ph.D.)

Nyitrai Konstantin Filozófus Egyetem (Szlovákia)

Hegedűs Orsolya (Ph.D.)

Nyitrai Konstantin Filozófus Egyetem (Szlovákia)

Németh Zsófia (Mgr.)

Czuczor Gergely Alapiskola (Szlovákia)

E-mail:

tszabo@ukf.sk

Cite: Idézés:

Szabó Tibor, Hegedűs Orsolya és Németh Zsófia (2025): Kreatív történetmesélés Ozobotok segítségével. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/4. szám. 99-108.

Doi: <https://www.doi.org/10.35406/MI.2025.2.99>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE:

Ethics Permission / Etikai engedély: KFS/2025/MI0013

Reviewers: Lektorok:

Public Reviewers / Nyilvános Lektorok:

1. Pšenáková Ildikó (Ph.D.), Nagyszombati Egyetem (Szlovákia)
2. Bakonyi Viktória (Ph.D.), Eötvös Loránd Tudományegyetem (Magyarország)

Anonymous reviewers / Anonim lektorok:

3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

A digitális technológia egyre nagyobb szerepet tölt be az oktatásban, különösen az élményszerű és kreatív tanulás támogatásában. Célunk, hogy megosszuk tapasztalatainkat az Ozobot robotok történetmesélésre való alkalmazhatóságáról az alapiskolás gyermekek körében. Kutatásunkat egy alapiskola negyedik osztályában végeztük. A gyermekek egy Ozobot robotokkal támogatott foglalkozássorozaton vettek részt, amelyet saját pályatervezés és történetmesélés követett. A foglalkozások szorosan kapcsolódtak az adott tantárgyak – matematika, természetrajz és honismeret – előírt tartalmához, amelyet a pedagógussal egyeztetettünk az érintett tanórák előtt. A tanulók előre megtervezett, rövid történetek köré épülő pályákat járattak be a robottal. Ahhoz, hogy a pálya teljesíthető legyen az Ozobot számára, részfeladatokat kellett megoldaniuk: a megfelelő színekkel kellett kiegészíteniük a hiányos, de tudatosan megtervezett pályákat. A helyes színezést követően a robot által vizuálisan is életre keltek a történet eseményei. A robot egyúttal módszertani segédeszközként is funkcionált. A foglalkozássorozatnak további célja az volt, hogy a tanulók megismerkedjenek az Ozobot robotokkal, mivel korábban még nem dolgoztak ilyen eszközzel. A záró foglalkozás során a gyermekek feladata egy saját pálya megtervezése és a hozzá kapcsolódó történet elmesélése volt a robot alkalmazásával. Tanulmányunkkal szeretnénk rámutatni arra, hogy az edukációs robotika és a történetmesélés ötvözése nemcsak műszaki készségeket hivatott fejleszteni, hanem komplex pedagógiai célok eléréséhez is hatékony eszközt kínál.

Kulcsszavak: digitális történetmesélés, Ozobot, kreativitás

Diszciplínák: pedagógia, informatika

Abstract*CREATIVE STORYTELLING WITH OZOBOTS*

Digital technology is playing an increasingly significant role in education, particularly in supporting experiential and creative learning. We aim to share our experience of using Ozobot robots for storytelling with primary school children. Our research was carried out in a fourth-grade class at a primary school. The children participated in a series of activities supported by Ozobot robots, followed by designing their own tracks and telling stories. The sessions closely aligned with the required content of the school subjects, as previously agreed with the teacher. We incorporated the activities into mathematics lessons twice and included single sessions in science and geography lessons. Students followed a pre-designed trajectory with the robot, based on a series of short stories. To make the track passable for the Ozobot, they had to complete sub-tasks: they needed to add the correct colour codes to the tracks, which were carefully designed but incomplete. After the correct colouring, the robot animated the story's events, effectively bringing the story itself to life. We can consider the robot a methodological

tool that has a significant motivational effect too. The series of activities also aimed to familiarise the children with Ozobot robots, as they had not worked with such devices before. In the last lesson the children's task was to create a track and tell a story with the help of the robot. Our study aims to demonstrate that educational robotics, when combined with storytelling, is not only intended to develop technical skills but also to support complex pedagogical objectives.

Keywords: digital storytelling, Ozobot, creativity

Disciplines: pedagogy, information technology

A digitális technológia térnyerése új lehetőségeket teremt az oktatásban, különösen az élményszerű és kreatív tanulás területén. Az edukációs robotika mint innovatív pedagógiai eszköz, egyre inkább előtérbe kerül, különösen az alsó tagozatos tanulók körében.

Kutatásunkat alapiskolás negyedik osztályos gyermekek körében valósítottuk meg. Szeretnénk kiemelni azt, hogy a gyermekekkel végzett foglalkozássorozatot a STEAM (Science, Technology, Engineering, Mathematics and Arts) modellen alapuló elvekkel összhangban végeztük. A modellnek kiemelt szerepe van az esélyegyenlőség előmozdításában. Sok lány azért nem választ műszaki vagy természettudományos pályát, mert az ezekhez szükséges ismereteket olyan módszerekkel tanították, amelyek nem igazodtak a lány tanulók sajátos tanulási igényeihez és képességfejlődéséhez (Babály és Kárpáti, 2016, Kárpáti és Pataiová, 2021). A robotok alkalmazása természetesen más célokat is szolgált, melyek közül az egyik legfontosabb a gyermekek motivációjának erősítése volt. Hiszen a tanár felkészültségén túl, sokat segíthet a megfelelő eszközök használata is. Pataiová és Kárpáti (2021) is felhívják a

figyelmet arra, hogy a motiváció megvalósítási szakaszában fontos, hogy a tanár magyarázata érdekes és világos legyen, illetve megfelelő módszereket és eszközöket használjon. Ezt az elvet mi is követni próbáltuk. Az egyik ilyen módszer a történetmesélés digitális formában.

Vivitsou (2020) szerint az oktatásban a digitális technológiák segítségével megvalósuló történetmesélés az iskolákban olyan hibrid környezetet teremt, amely egyaránt támogatja a tanulók és a pedagógusok cselekvését és értelmezési folyamatait. A digitális technológia eszközei nagyon gyakran valamilyen edukációs robotok formájában jelennek meg, mint például a BeeBot vagy az Ozobot.

Az informatikai problémák dialógusos feldolgozásának és a történetek vizuális formában történő kreatív ábrázolásának kombinációja tartósan elősegítheti a gyermekek narratív nyelvhasználatának fejlődését (Tengler et al., 2022). Lanszki (2019) kutatása is arra utal, hogy a digitális történetmesélés fejleszti a kreativitást és a narratív gondolkodást.

Jelen tanulmányunk célja, hogy bemutassa az Ozobot robot különböző tantárgyakba való integrációjának lehetőségeit, valamint feltárja, hogy a gyermekek milyen módon jutnak el a

saját digitális történetmesélés és pályaalkotás szintjére. Mindez hatással van például a kreativitásukra is.

Módszerek

A kutatásban résztvevő gyermekek egy Ozobot robotokkal támogatott foglalkozássorozaton vettek részt, amelyet saját pálya-tervezés és történetmesélés zárt. Az utolsó foglalkozáson a tanulók párban dolgoztak – összesen 10 tanulópár, ellentétben a korábbi foglalkozásokkal, ahol egyéni munkában tevékenykedtek. A foglalkozásokat hagyományos osztályteremi környezetben valósítottuk meg.

Módszertani segédeszközként az Ozobot Bit+ típusú kis méretű robotokat használtuk. Ezek a robotok speciális filctollal rajzolt vonalak mentén közlekednek. Amikor a pályán egy színkóddal megszakított szakaszhoz érnek, felismerik a színkombinációt, majd végrehajtják a gyártó által előre definiált utasítást. A robot működésének részletes leírását többek között Szabó és Jakab (2024) tanulmánya is ismerteti, továbbá a robotot gyártó cég hivatalos honlapján is megtalálható (URL: <https://ozobot.com>).

Foglalkozássorozat

Célunk elérése érdekében – tehát hogy a tanulók képesek legyenek saját történetet alkotni, és azt Ozobot robot segítségével bemutatni –, több előkészítő foglalkozáson is részt kellett venniük. Összesen négy tanóra, ill. ezáltal négy foglalkozás kidolgozására és lebonyolítására került sor az adott csoporttal

(4. osztály). A tevékenységek három tantárgy keretében valósultak meg: két matematika órán, valamint egy-egy honismeret és természetismeret tanórán. Ezekből kettőt részletebben is bemutatunk. Fontos hangsúlyoznunk, hogy a kidolgozott foglalkozások összhangban álltak az adott tantárgyak előírt tartalmi követelményeivel, emellett lehetőséget teremtettek a tantárgyközi kapcsolatok megjelenítésére is. Úgy véljük, hogy ez nem ritka jelenség az edukációs robotika alkalmazása során (Szabóné, 2023; Mező és Szabóné, 2021).

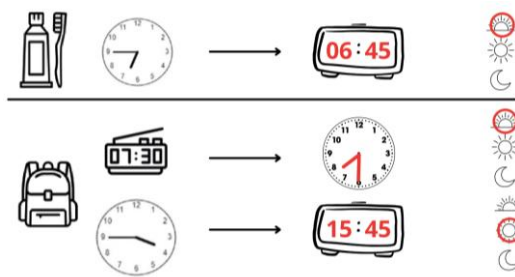
Matematika – Napirend

Az Ozobot Bit+ robotokkal végzett tevékenység célja az volt, hogy a tanulók játékos formában ismerkedjenek meg a mindennapi teendők és feladatok sorrendiségével és időrendjével. A foglalkozás középpontjában egy napirend-pálya állt, amely egy gyermek – jelen esetben az Ozobot robot által megjelenített szereplő – napi programját mutatta be.

A pálya tervezése során a fő szempont az volt, hogy a diákok számára ismerős, hétköznapi események (például ébredés, reggeli, iskolába indulás, tanulás, játék, vacsora, lefekvés) logikus sorrendben kövessék egymást. A robot mozgása így módon nemcsak a programozási logika szemléltetését szolgálta, hanem hozzájárult a napi rutinok tudatosításához is.

A pálya olyan helyeket tartalmazott a színkódok kitöltésére, amelyeknél segítséget is feltüntettünk (lásd: 1. ábra). Azokon a helye-

2. ábra: A napirend foglalkozásból tartozó feladatlap – részlet. Forrás: a szerzők



A pálya története a következőképpen épült fel: reggel a diákot jelképező robot felébred, felgyorsul, és siet az iskolába. Indulás előtt, a reggeli rutinhoz tartozó állomásoknál (például fogmosás, reggeli) a robot megáll, majd az iskolába érkezés után lelassul. Délután zongoraóra és tanulás következik, ekkor a sebesség újra mérséklődik. Az esti olvasás során a robot

a leglassabb sebességfokozatra vált, jelezve a nyugodtabb, pihenésre hangoló időszakot. Vacsorakor ismét megáll, végül az alvás állomáshoz érve újra felgyorsul — hiszen miután „elalszik”, a történet szerint gyorsan ismét reggel lesz.

Honismeret – Érsekújvári vár

Már az elején le kell szögeznünk, hogy az itt ismertetett gondolatmenet számos más helyszínre és történelmi eseményre is adaptálható. Az érsekújvári vár választása számunkra kézenfekvő volt, mivel a vizsgálatban érintett iskola is ebben a városban található. Bár a vár napjainkra már nem maradt fenn, emlékét őrzik a bástyák nevei, amelyeket ma az egyes utcák viselnek.

A foglalkozást azzal indítottuk, hogy kiemeltük a vár jelentős szerepét a török időkben. Ismertettük a legfontosabb történelmi eseményeket: meddig állt ellen a török erőknél, mikor esett el és mikor szabadult fel. Mindeközben kivetítettük a vár történelmi alaprajzát. Ezt követően rákérdeztünk arra, hogy a gyerekek tudják-e, hol található az egyes utcák, amelyek a bástyákról voltak elnevezve, hiszen ezek jelzik a bástyák egykori elhelyezkedését. A történetbe további objektumokat (például nevezetes épületeket) is bevonhatunk, amelyek segíthetnek a térbeli tájékozódásban. Mutathatunk képeket az említett utcákról vagy a volt bástyák közelében álló épületekről, hogy a gyerekek könnyebben be tudják azonosítani a tárgyalt helyszíneket. A tájékozódást tovább segítheti az égtájak

megismertetése, melyek alapján pontosabban be lehet határolni az egyes bástyák egykori helyét a vár szerkezetén belül. Például a következőképpen:

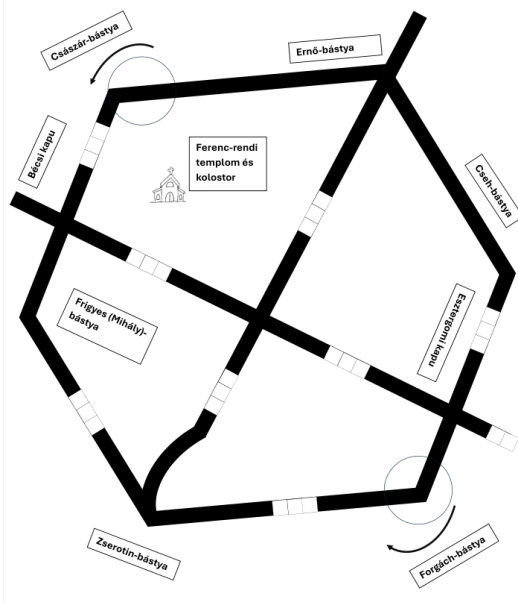
- Ernő-bástya: É – Zámčan, Takarékpénztár,
- Cseh-bástya: ÉK – zsinagóga,

és így tovább.

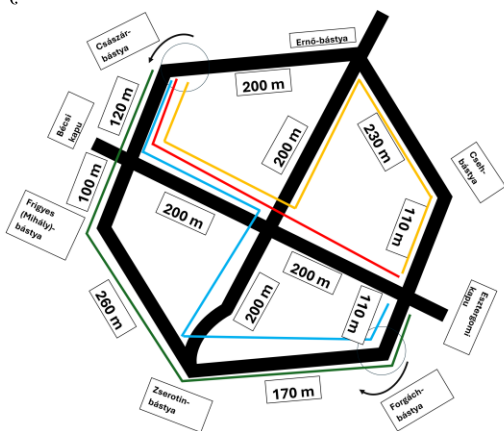
Mindezek után került sor a történetünk főszereplőjeként megjelenő Ozobot robot feladatára. Ozo visszarepült az érsekújvári vár fénykorába, amikor a környéket gyakran kellett megvédeni az oszmán portyázóktól. Ozo maga is kivette a részét a harcokból. Egy alkalommal éppen a Császár-bástyánál tartózkodott, miközben az Esztergomi kapu környékén egy kisebb portyázó csapat fosztogatta az ott élő lakosságot. Parancsot kapott arra, hogy csapatával űzze el a támadókat, ehhez azonban el kellett jutnia az Esztergomi kapuhoz. A feladat tehát az volt, hogy a tanulók segítsék Ozót elérni a célpontot. A történet ismertetését követően a diákok megkapták a vár alaprajza szerint elkészített pályát (lásd 3. ábra).

A pályán kijelöltük a színkódok helyét, azonban nem tüntettük fel, mely színkombinációkat kell alkalmazniuk. A feladat során négyféle színkód használatára volt lehetőség: a célpontnál a játék végét jelző kódra (amelynek hatására a robot „táncol”, „ünnepel”), míg a pálya különböző kereszteződéseiben három további színkombinációra, amelyek jobbra, balra vagy egyenesen irányítják a robotot. Kivetítettük a 4. ábrán látható térképet, majd közösen megjelöltünk több lehetséges útvonalat, amelyek mentén Ozo eljuthat az Esztergomi kapuhoz.

3. ábra: A hatszög a ma már nem létező vár alaprajzát reprezentálja, amelynek csúcspontjai a korabeli bástyák elhelyezkedését jelölik. Forrás: a szerzők



4. ábra: A hatszögalapú vár szomszédos bástyáinak távolsága egymástól, ill. a főtér távolsága az egyes bástyáktól vagy kapuktól. Színes törtvonalakkal az egyes lehetséges útvonalakat jelöljük. Forrás: a szerzők



Fontos volt, hogy ezt minél gyorsabban megtegye, így a tanulóknak a legrövidebb útvonalat kellett meghatározniuk. Az ideális eset megtalálásához az egyes útvonalak hosszát a térképről leolvasható útszakaszok összeadásával számították ki. Ezek a szakaszhosszak valósághoz közeli, azonban az egyszerűbb kezelhetőség érdekében tízesekre kerekített értékeket tüntettünk fel. A számítások helyességét a 1. táblázat kivetítésével közösen ellenőriztük.

1. táblázat: Az útvonalak hosszának kiszámítása (a táblázat sorainak színe a megfelelő útvonalak színével van összhangban). Forrás: a szerzők

Útvonalak	Távolság [méter]
1. útvonal	$120 + 100 + 260 + 170 + 110 = 760$
2. útvonal	$120 + 200 + 200 = 520$
3. útvonal	$120 + 200 + 200 + 170 + 110 = 800$
4. útvonal	$120 + 200 + 200 + 230 + 110 = 860$

Tekintettel arra, hogy a 4. évfolyamos tanulók már magabiztosan végeznek összeadást 1000-es számkörben, valamennyi kapott eredmény ebben a tartományban maradt. A legrövidebb útvonal kiválasztását követően a tanulók a megfelelő színek kódokat alkalmazva irányították a robotot a kapu irányába. Olyan tanuló is akadt, aki pusztán vizuális becslés alapján is képes volt meghatározni a leg-

rövidebb útvonalat; ennek ellenére megkértük, hogy számításokkal is támassza alá az állítást.

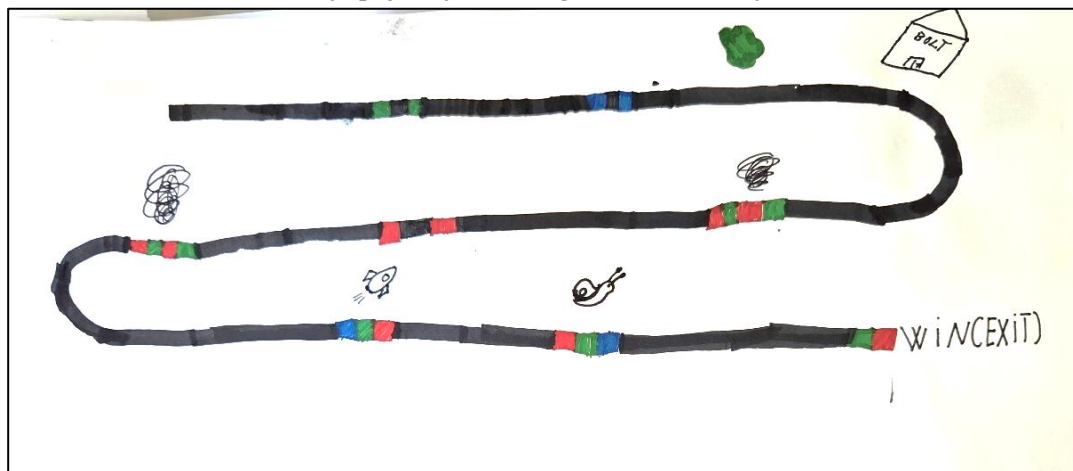
A foglalkozás zárásaként a tanulók kaptak még egy egyszerűbb feladatot. Ebben Ozo egy másik napon a Forgách-bátyán teljesített szolgálatot, azonban új parancsot kapott: ellenőriznie kellett a Zserotin-bátyát, majd át kellett mennie őrködni az Ernő bátyára. Ez azt jelenti, hogy egy köztes bátyán áthaladva jutott el a célállomásra. Ebben az esetben már csupán egyetlen lehetséges útvonal állt rendelkezésére, hiszen az előző részfaladat során a gyerekek több, a színek kódok számára kihagyott mezőt is felhasználtak.

Eredmények: a tanulók számára egy tanóra állt rendelkezésre saját pályájuk megtervezésére, és a hozzá kapcsolódó történet kidolgozására. Ezt követően a robot segítségével bemutatták alkotásukat osztálytársaiknak. A tervezési fázisban lehetőségük nyílt teljesen

üres lapra, önállóan dolgozni, és a pályát tetszőleges módon kiegészíteni, például különböző rajzokkal (lásd 5. ábra).

A párban végzett kooperatív munka a legtöbb esetben megvalósult: a tanulók közösen fogalmazták meg a történetet, együtt készítették el a pályát, és egyeztettek egymással az elképzeléseiket. Egy tanulópáros kivételével minden gyerek motiváltan vett részt a foglalkozáson, amit az is bizonyított, hogy szívesen mutatták be az elkészült pályáikat, és mesélték el történeteiket. A pálya tervezése minden esetben lineáris történetvezetésre épült, azaz, egyik pálya sem tartalmazott elágazást, holott korábbi feladatokban a tanulók már találkoztak ilyen lehetőséggel. A leggyakrabban előforduló nehézséget a színek kódok megfelelő megfestése, a túl sötét árnyalatok vagy az aránytalan nagyságú részek színkódon belüli használata jelentette.

5. ábra: A tanulók által tervezett pálya az általuk megalkotott történethez. Forrás: tanulók alkotása



Pozitív eredményként említhető, hogy a történetek illeszkedtek a robot mozgásához, és mindegyik narratívának volt jól elkülöníthető eleje, közepe és vége. A robot vizuális visszajelzései tovább erősítették a történet élményszerűségét és szemléletességét. Megfigyeléseink alapján a tanulók többsége képes volt strukturált, koherens történetet létrehozni, azt vizuális elemekkel kiegészíteni, és összekapcsolni az Ozobot mozgásával.

Befejezés

Összegzésként megállapítható, hogy a tanulók aktívan, kreatívan és motiváltan vettek részt a feladatok elvégzésében. A történetmesélés és a robotika kombinációja komplex – a problémamegoldó gondolkodás, a kommunikáció és a kooperáció területén egyaránt jelentős – készségfejlesztést tesz lehetővé.

Meggyőződésünk, hogy a történetmesélés mint kreatív tanulási forma kiválóan alkalmas arra, hogy a technológiai eszközöket élményszerűen és motiváló módon kapcsolja be az oktatási folyamatba, és integrálja a tananyagba.

Tanulmányunk célja az volt, hogy rámutassunk: az edukációs robotika nem csak digitális kompetenciák fejlesztésére alkalmas, hanem komplex pedagógiai célokat is támogat.

Irodalom

Babály, B. és Kárpáti A. (2016). The impact of creative construction tasks on visuospatial information processing and problem solving. *Acta Politechnica Hungarica*, 13 (7), pp. 159-180. Doi:

<https://www.doi.org/10.12700/APH.13.7.2016.7.9>

Kárpáti, A. és Pataiová, H. (2021).

Természettudományos és művészeti integráció 2: A STEAM modell a gyakorlatban. *Katedra*, 29 (4), pp. 26-28.

Lanszki, A. (2019). Tanulói kreativitás fejlesztése digitális történetmesélés segítségével. *Iskolakultúra*, 29 (4-5), pp. 71-85. Doi:

<https://www.doi.org/10.14232/ISKKULT.2019.4.5.71>

Mező, K. és Szabóné Burik, E. (2021): A robotokkal történő oktatás, az élménypedagógia aspektusából. *Mesterséges intelligencia – interdiszciplináris folyóirat*, III. évf., 2021/2. szám. 19-32. Doi:

<https://www.doi.org/10.35406/MI.2021.2.19>

Pataiová, H. és Kárpáti, A. (2021). A motiváció jelentősége az edukációs folyamatban. *Katedra*, 29 (2), pp. 24-27.

Szabó, T. és Jakab, Zs. (2024). Didaktikai játékok a matematika oktatásában Ozobot robot felhasználásával. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, 6 (1). pp. 87-103. Doi:

<https://www.doi.org/10.35406/MI.2024.1.87>

Szabóné Balogh Ágota (2023): Mesterséges intelligencia az oktatásban. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, V. évf. 2023/2. szám. 51-61. Doi:

<https://www.doi.org/10.35406/MI.2023.2.51>

Tengler, K., Kastner-Hauler, O., Sabitzer, B. és Lavicza, Zs. (2022). The Effect of

Robotics-Based Storytelling Activities on Primary School Students' Computational Thinking. *Education Sciences*, 12 (1), pp. 1-15. Doi: <https://www.doi.org/10.3390/educsci12010010>

Vivitsou, M. (2020). Digital Storytelling in Teaching and Research. In Tatnall, A. (Eds.). *Encyclopedia of Education and Information Technologies*, pp. 573-588. Doi: https://www.doi.org/10.1007/978-3-030-10576-1_58

**GITÁRZENE A MESTERSÉGES INTELLIGENCIA KORÁBAN
– GONDOLATOK A XV. GYÖNGYÖSI NEMZETKÖZI GITÁRFESZTIVÁL
KAPCSÁN**

Szerző:

Papp Sándor (Ph.D.)
Miskolci Egyetem
(Magyarország)

Mező Kristóf Szíriusz
Kocka Kör
(Magyarország)

E-mail:

sziriusz.mezo1@gmail.com

**Cite:
Idézés:**

Papp Sándor és Mező Kristóf Szíriusz (2025): Gitárzene a mesterséges intelligencia korában – Gondolatok a XV. Gyöngyösi Nemzetközi Fitarfesztyál kapcsán. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/4. szám. 109-115.

Doi: <https://www.doi.org/10.35406/MI.2025.2.109>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE:

Ethics Permission / Etikai engedély: KFS/2025/MI0014

**Reviewers:
Lektorok:**

Public Reviewers / Nyilvános Lektorok:

1. Ritter József (Ph.D.), Miskolci Egyetem (Magyarország)
2. Széplaki Zoltán (Ph.D.), Miskolci Egyetem (Magyarország)

Anonymous reviewers / Anonim lektorok:

3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

XV. Gyöngyösi Nemzetközi Gitár-fesztivál 2025. július 24-27. között került megrendezésre Gyöngyösön, Magyarországon. E zenei esemény apropójából jelen tanulmány inspiráló, témafelvető jelleggel hívja fel a figyelmet a gitárzene és a mesterséges intelligencia kapcsolódási pontjaira.

Kulcsszavak: gitár, zene, MI

Diszciplínák: pedagógia, zene

Abstract

GUITAR MUSIC IN THE AGE OF ARTIFICIAL INTELLIGENCE – THOUGHTS IN CONNECTION TO THE 15TH GYÖNGYÖS INTERNATIONAL GUITAR FESTIVAL

The 15th Gyöngyös International Guitar Festival was held in Gyöngyös, Hungary, between 24 and 27 July 2025. On the occasion of this musical event, the present study adopts an exploratory and inspirational approach, drawing attention to the points of intersection between guitar music and artificial intelligence.

Keywords: guitar, music, artificial intelligence

Disciplines: pedagogy, music

A „XV. Gyöngyösi Nemzetközi Gitár-fesztivál megvalósítása” címmel 2025-ben valósított meg művészeti tehetséggondozást (is) népszerűsítő projektet a Kocka Kör Tehetséggondozó Kulturális Egyesület (lásd: 1. ábra).

A rendezvény 2025. január 1. és 2025. április 10. között a Nemzeti Kulturális Alap (www.nka.hu) Zeneművészet Kollégiuma támogatásával kezdett megvalósulni (pályázati azonosító: 110124/ 00855). A pályázatot a Nemzeti Kulturális Alapról szóló 1993. évi XXIII. törvény végrehajtásáról szóló 9/2006. (V. 9.) NKÖM rendelet 1. számú mellékletének értelmében 2025. április 11. napjától

az Előadó-művészetek Kollégiuma vette át (a projekt új azonosítószáma 502124/1754 lett).

A XY. Gyöngyösi Nemzetközi Gitárfesztivált továbbá Gyöngyös Város Önkormányzata és a gyöngyösi Vachott Sándor Városi Könyvtár is segítette (2. ábra).

A XV. Nemzetközi Gitárfesztivál-sorozat alapító-főszervezője: Dr. Papp Sándor, a Miskolci Egyetem Zeneművészeti Kar dékánja, aki a rendezvénysorozat kezdete óta igazgatja és szervezi az események megvalósulását, s közben a hazai és külföldi gitárművészek szakmai kapcsolatteremtésének lehetőségeit, valamint a gitárzene népszerűsítését a nagyközönség körében.

1. ábra: A XV. Gyöngyösi Nemzetközi Gitárfesztivál szervező civil szervezet



2. ábra: A XV. Gyöngyösi Nemzetközi Gitárfesztivál támogatói és fellépői



Fellépők:

2025.07.24.: Juan Lorenzo
2025.07.25.: Papp Sándor és Mező Kristóf Színusz
2025.07.26.: Enyedi Sándor
2025.07.27.: Nagy Tibor (Wyrág) és Szabó Dénes

Részletek: www.kockakor.hu



SZERVEZŐ:

 **KOCKA KÖR**
www.kockakor.hu

VEDNÖK:
Dr. Papp Sándor
dékán
Miskolci Egyetem
Bartók Béla
Zeneművészeti Kar



TÁMOGATÓK:


**GYÖNGYÖS VÁROS
ÖNKORMÁNYZATA**


**VACHOTTI
SÁNDOR
VÁROSI
KÖNYVTÁR
GYÖNGYÖS**

 **Nemzeti
Kulturális
Alap**

A rendezvényt a
Nemzeti Kulturális Alap
Előadó-művészetek
Kollégiuma támogatja.
Pályázati azonosító:
502124/1754

Gitárzene és mesterséges intelligencia

A zenében napjainkban új színfoltot jelent a mesterséges intelligencia (MI) zenei komponálást, -szerkesztést és közzétételt részben vagy egészben megvalósító megjelenése (lásd például: Briot, Hadjeres és Pacht, 2020, Engel és tsai., 2017, Huang és tsai, 2018, Humphrey, Bello és LeCun, 2013, McFee és tsai, 2015, Roads, 2015). Nincs ez másként az klasszikus, akusztikus, illetve elektromos gitárral megvalósuló zeneszerzés, -előadás esetében sem.

Az alábbiakban mesterséges intelligencia és a gitárzene kapcsolódási pontjai mentén kerül bemutatásra egy lista néhány lehetséges kutatási témáról. A lista előkészítése során a „gitárzene a mesterséges intelligencia korában” témakörhöz kapcsolódó tematikus szempontok rendszerezéséhez a ChatGPT (OpenAI) generatív mesterséges intelligencia rendszere került felhasználásra vázlatkészítő és gondolatgeneráló eszközként. Megjegyzendő, hogy ez már önmagában is egy mesterséges intelligencia alapú online platform gitárzenére fókuszáló egyik felhasználási lehetőségét demonstrálja. Lehetséges témák és összefüggések tehát az alkalmazott mesterséges intelligencia felvetései alapján a gitárzene és a mesterséges intelligencia témakörök kapcsán (3. ábra):

1. A mesterséges intelligencia szerepe a gitárzenében

- MI mint komponáló eszköz (dallam-, harmónia-, riff-generálás)
- MI mint társalkotó (ember–AI együttműködés)
- MI mint elemző rendszer (stílus-, játék-technika-elemzés)

2. Gitárhang és hangszín modellezése

- Gitárerősítők és effektek neurális modellezése (például amp modeling)
- Analóg hangzás digitális rekonstruálása deep learninggel
- Fizikai modellezés vs. adatvezérelt MI-megoldások

3. Gitártechnika és játékmód MI-alapú elemzése

- Pengetési technikák (váltott, sweep, fingerstyle) felismerése
- Expresszív elemek (bend, vibrato, slide) automatikus detektálása
- Előadói stílusok (például: blues, metal, jazz) gépi osztályozása

4. MI-alapú zeneszerzés gitárra

- Riff- és akkordmenet-generálás neurális hálókkel
- Gitárcentrikus generatív modellek korlátai
- Stílustranszfer (például: „klasszikus gitár jazz stílusban”)

5. Gitározás tanulása és oktatása az MI korában

- Automatikus visszajelzés (intonáció, ritmus, artikuláció)
- Adaptív gyakorlórendszerek gitárosoknak
- MI-alapú virtuális tanár vs. emberi pedagógus

6. Improvizáció és valós idejű MI-rendszerek

- Valós idejű kísérogitár vagy backing track generálás
- Interaktív improvizáció ember és MI között
- Reagáló rendszerek (tempó, hangnem, dinamika követése)

7. Kreativitás, szerzőség és autenticitás

- Ki a szerző: a gitáros, a fejlesztő vagy az algoritmus?
- Az „emberi kéz” szerepe a digitális gitárzenében
- Az autentikus gitárhang jelentésének átalakulása

8. Etikai és jogi kérdések

- Tréningadatok kérdése (gitárosok stílusának felhasználása)
- Szerzői jog MI által generált gitárzenénél
- Előadói jogok és hangzásutánpótlás

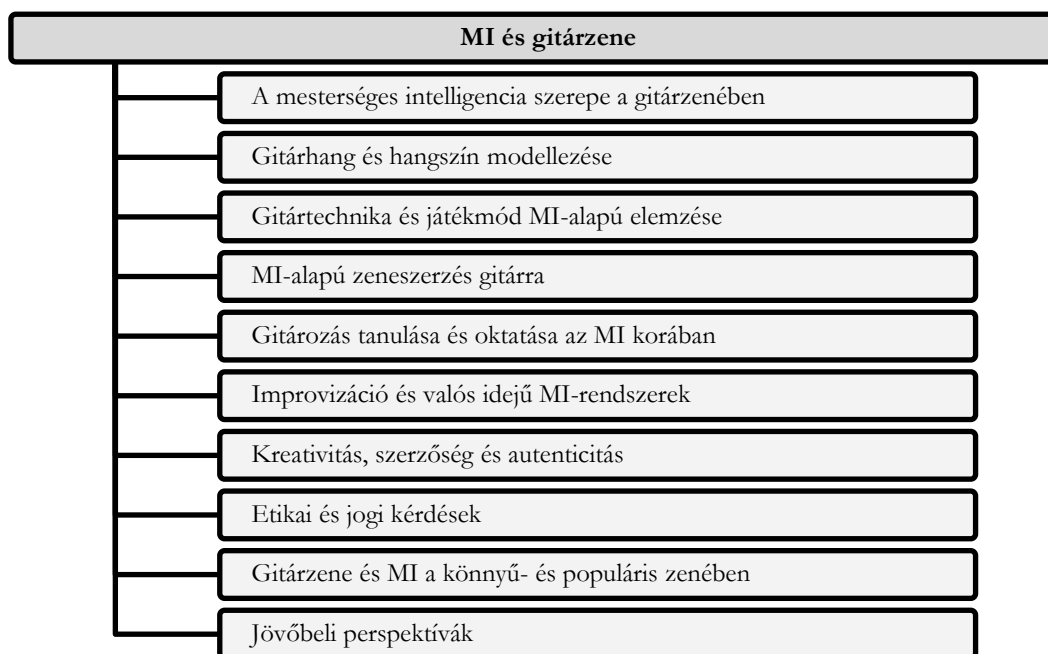
9. Gitárzene és MI a könnyű- és populáris zenében

- MI-alapú gitársávok a pop- és filmzenében
- Produceri döntések automatizálása
- „Session gitáros” szerepének változása

10. Jövőbeli perspektívák

- Teljesen virtuális MI-gitárosok
- Hibrid hangszerek (fizikai gitár + MI-réteg)
- Gitár mint interfész az MI-alapú zeneszerzéshez

3. ábra: A mesterséges intelligencia (MI) téma és gitárzene összefüggései, mint lehetséges kutatási témakörök, témafelvetések. Forrás: a Szerző



Zárógondolatok

Az első (úgynevezett Torres-féle formával jellemezhető) klasszikus gitár 1850-ben jelent meg (Wade, 2010), az elektronikus gitár „születése” pedig nyolc évtizeddel későbbre, 1931-re datálható (Millard, 2004). A napjainkig eltelt 94-175 évben mind a klasszikus, mind az elektromos gitárzene nagy népszerűségnek örvendő, sokoldalú hangzást biztosító, egyedi zenei élményvilágot nyújtani képes művészeti területté nőtte ki magát. Ez azonban nem jelenti azt, hogy ne lehetne összefüggésbe hozni a napjainkban egyre népszerűbb, bár (a nagyközönség által online elérhető formában) alig néhány éves múlttal rendelkező mesterséges intelligencia technológiával (a generatív mesterséges intelligencia technológia a 2020-as évektől elérhető, a jelen tanulmány írásához felhasznált ChatGPT első közreadott verziója 2022-ben jelent meg, s jelen tanulmány írásakor mindössze három éves múlttal rendelkezik).

Köszönetnyilvánítás

Köszönjük a Nemzeti Kulturális Alap Előadó-művészetek Kollégiuma által kezelt pályázati támogatást (pályázati azonosítószám: 502124/1754).

Irodalom

Briot, J.-P., Hadjeres, G., & Pachet, F. (2020). *Deep learning techniques for music generation*. Springer. Doi:

<https://doi.org/10.1007/978-3-319-70163-9>

Engel, J.; Resnick, C.; Roberts, A.; Dieleman, S.; Eck, D.; Simonyan, K.; & Norouzi, M. (2017). *Neural audio synthesis of musical notes with WaveNet autoencoders*. Proceedings of the 34th International Conference on Machine Learning. Doi:

<https://doi.org/10.48550/arXiv.1704.01279>

Huang, Cheng-Zhi Anna; Vaswani, Ashish; Uszkoreit, Jakob; Shazeer, Noam; Simon, Ian; Hawthorne, Curtis; Dai, Andrew M.; Hoffman, Matthew D.; Dinculescu, Monica; & Eck, Douglas (2018). *Music transformer: Generating music with long-term structure*. International Conference on Learning Representations (ICLR). Doi:

<https://doi.org/10.48550/arXiv.1809.04281>

Humphrey, E. J., Bello, J. P., & LeCun, Y. (2013). Feature learning and deep architectures: New directions for music informatics. *Journal of Intelligent Information Systems*, 41(3), 461–481. Doi:

<https://doi.org/10.1007/s10844-013-0248-5>

McFee, B.; Raffel, C.; Liang, D.; Ellis, D.P.W.; McVicar, M.; Battenberg, E. & Nieto, O. (2015). *librosa: Audio and music signal analysis in Python*. Proceedings of the 14th Python in Science Conference, 18–25. Doi:

<https://doi.org/10.25080/Majora-7b98e3ed-003>

- Millard, A.J. (2004). *The Electric Guitar: A History of an American Icon*. Johns Hopkins University Press
- OpenAI. (2023). *ChatGPT* (GPT-4 version) [Large language model]. Megnyitva: 2025.12.01. URL: <https://chat.openai.com/>
- Roads, C. (2015). *Composing electronic music: A new aesthetic*. New York: Oxford University Press.
- Wade, G. (2010). *A concise history of the classic guitar*. Mel Bay Publications.

INTEGRATING DIGITAL PEDAGOGY INTO TRADITIONAL MUSIC EDUCATION

Author(s) / Szerző(k):

Loredana Muntean (Ph.D.)
University of Oradea (Romania)

Morel Koren (Ph.D.)
Bar Ilan University (Israel)

E-mail:

mt.loredana@gmail.com

Cite: Muntean, Loredana and Koren, Morel (2025): Integrating Digital Pedagogy
Idézés: into Traditional Music Education. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 117-125.
Doi: <https://www.doi.org/10.35406/MI.2025.2.117>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE: Ethics Permission / Etikai engedély: KFS/2025/MI0015

Reviewers: Anonymous reviewers / Anonim lektorok:

- Lektorok:**
1. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 2. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
 4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Abstract

Traditional approaches to music pedagogy, rooted in the work of eminent twentieth-century educators (e.g., Kodály, Orff), emphasize the cultivation of musical sensitivity through singing, solmization, rhythmic exercises, and improvisation. These practices remain central in

contemporary music education. At the same time, digital pedagogy introduces a new dimension, facilitating the design, delivery, and evaluation of instruction through interactive platforms, learning management tools, and virtual learning spaces. This study focuses on the pedagogical integration of artificial intelligence (AI)-based software within the context of general music education. The focus is on the Solfy program, designed both as an interactive classroom tool and as support for independent home practice. Its distinctive feature is the AI-driven voice recognition system that delivers immediate, personalized feedback, enabling differentiated development of vocal skills. The study shows that Solfy can effectively complement traditional pedagogical methods at the pre-vocational stage. By integrating AI-based tools into established practices, music education can become more efficient, accessible, and reflective, for example, in areas such as intonation accuracy, tonal stability, and vocal confidence. Immediate feedback, error detection, and reward mechanisms (e.g., orchestral accompaniment) can foster consistent practice, reduce performance anxiety, and enhance learner autonomy by encouraging independent, goal-oriented learning strategies.

Keywords: Environmental noise; MATLAB; tensor data; singular value decomposition; smart city monitoring

Discipline: Environmental Engineering

Absztrakt

A DIGITÁLIS PEDAGÓGIA INTEGRÁLÁSA

A HAGYOMÁNYOS ZENEOKTATÁSBAN

A zenepedagógia hagyományos megközelítései, amelyek kiemelkedő huszadik századi pedagógusok (pl. Kodály, Orff) munkásságában gyökereznek, a zenei érzékenység fejlesztését hangsúlyozzák éneklés, szolmizáció, ritmikus gyakorlatok és improvizáció révén. Ezek a gyakorlatok továbbra is központi szerepet játszanak a kortárs zenei nevelésben. Ugyanakkor a digitális pedagógia új dimenziót vezet be, elősegítve az oktatás tervezését, lebonyolítását és értékelését interaktív platformok, tanulásmenedzsment eszközök és virtuális tanulási terek segítségével. Ez a tanulmány a mesterséges intelligencia (MI) alapú szoftverek pedagógiai integrációjára összpontosít az általános zenei nevelés kontextusában. A hangsúly a Solfy programon van, amelyet interaktív tantermi eszközként és az önálló otthoni gyakorlás támogatására is terveztek. Megkülönböztető jellemzője a mesterséges intelligencia által vezérelt hangfelismerő rendszer, amely azonnali, személyre szabott visszajelzést ad, lehetővé téve a vokális készségek differenciált fejlesztését. A tanulmány kimutatja, hogy a Solfy hatékonyan kiegészítheti a hagyományos pedagógiai módszereket a szakmai előkészítő szakaszban. A mesterséges intelligencia alapú eszközök bevett gyakorlatokba való integrálásával a zenei nevelés hatékonyabbá, hozzáférhetőbbé és reflektívabbá válhat például olyan területeken, mint az

intonációs pontosság, a hangszíntabilitás és az énekhanggal kapcsolatos magabiztosság. Az azonnali visszajelzés, a hibaészlelés és a jutalmazási mechanizmusok (pl. zenekari kíséret) elősegíthetik a következetes gyakorlást, csökkenthetik a teljesítmény miatti szorongást és fokozhatják a tanuló autonómiáját azáltal, hogy ösztönzik a független, célorientált tanulási stratégiákat.

Kulcsszavak: környezeti zaj; MATLAB; tenzoradatok; szingulárisérték-felbontás; okosváros-monitorozás

Diszciplína: környezetmérnöki tudomány

Digital pedagogy refers to the processes of designing, implementing, and evaluating learning supported by digital technologies. It encompasses educational platforms, interactive resources, and intelligent applications capable of providing immediate feedback and adapting tasks to learners' individual needs (European Commission, 2021; UNESCO, 2023; Zhang & Zhang, 2024). In recent years, artificial intelligence (AI) has emerged as a particularly promising tool for the personalization of learning experiences, contributing simultaneously to the development of musical and digital competences (Li et al., 2025; Han, 2025).

Traditional music education—grounded in vocal training, solfège, improvisation, and instrumental performance, and shaped by paradigms such as the Kodály and Orff approaches—remains an essential methodological reference point. The introduction of digital technologies does not displace these practices but rather extends and amplifies them, creating a productive dialogue between tradition and innovation. In this context, technology does not substitute the teacher;

instead, it strengthens the teacher's role as mediator of the musical experience.

Current scholarship identifies three major directions for applying AI in music education. The first concerns feedback and self-monitoring, where applications such as Solfy or SmartMusic enable immediate correction of intonation and rhythm, fostering autonomy and efficiency in individual practice (Li et al., 2025; Nichols, 2014). The second involves AI-assisted creativity and expressivity, through generative tools that support adolescents' creative activity and reduce performance-related anxiety (Hu et al., 2025; Cheng, 2025). The third focuses on algorithmic models for teaching and assessment, developed mainly in higher education, which employ neural networks or large language model-based agents to support music theory learning and the analysis of vocal performance (Han, 2025; Jin et al., 2025; Wei et al., 2022).

In the Romanian context, Pop-Sârb (2021) demonstrated that integrating digital solutions and AI-based applications can provide valuable pedagogical support for renewing

school music education. At the same time, Popean (2022) warned against the risk of absolutizing technology and losing the authentic artistic dimension in the absence of critical reflection. More recent contributions emphasize the role of singing and (self-)evaluation in broad-based music literacy, as evidenced by Muntean, Weidenfeld, and Koren (2022), as well as in research presented at ISME (Muntean & Koren, 2022). In this regard, applications such as Solfy operationalize the principles of music literacy through singing and self-assessment in an accessible format that combines progress monitoring, individualized development, and reflective feedback.

International analyses confirm that AI has the potential to accelerate learning, personalize educational trajectories, and foster reflective dimensions of music training—provided that it is implemented responsibly and supported by ongoing teacher development (Merchán Sánchez-Jara et al., 2024; O’Leary, 2025). Thus, the integration of digital pedagogy and AI tools should be understood not simply as a technological innovation but as a methodological transformation, one in which educational traditions are interwoven with emerging practices, preparing learners for music education attuned to the realities of the twenty-first century.

Against this conceptual and empirical background, the present study advances the discussion by examining the use of Solfy at the pre-vocational level—an educational stage less frequently explored in the international literature—thereby contributing new insights

into how AI can support early vocal development while complementing traditional pedagogical methods.

Unlike previous studies that primarily examined AI integration at university or advanced levels, the present study focuses on the pre-vocational stage, contributing new empirical insights into how AI can support early vocal development.

Demonstration of using Solfy in a classic musical development

The example below demonstrates the possibility of integrating the Solfy (an AI-based digital tool) into individual vocal training. This case shows not only the participant’s technical progress but also the motivational and self-reflective dynamics associated with learning by the example of a 14-year-old student preparing for admission to a vocational high school with a music specialization. Ethical considerations were carefully observed: informed consent was obtained from the student’s parents, and the participant voluntarily agreed to take part in the diagnostics and development. The research was conducted in line with institutional ethical guidelines for studies involving minors. Written parental consent and participant assent were obtained prior to data collection, and all data were anonymized.

At baseline, her competence level was characterized by unstable intonation, limited breath control, and reduced experience in repertoire performance.

The intervention extended over three months (March–May 2024) and was struc-

tured into two successive stages. The first stage involved established traditional methods, including breathing exercises, intonation and solfège training through the Kodály method, and rhythmic activities inspired by Dalcroze eurhythmics. The second stage introduced the Solfy application, an AI-based educational tool that provides instant feedback on intonation and rhythm, monitors error typologies, and offers motivational reinforcement through orchestral accompaniment.

Data were collected using a combination of complementary methods:

- direct observation of practice sessions;
- audio recordings taken at the beginning and end of the intervention;
- Solfy-generated reports documenting the frequency and typology of errors;
- the participant's reflective journal, where she recorded perceptions of progress and motivation.

The analysis combined qualitative and quantitative procedures. Vocal performances were comparatively assessed at baseline and post-intervention; intonation errors and tonal stability were quantified using Solfy's analytical reports; and information derived from observation and journaling was triangulated to strengthen interpretive validity. Pitch stability was assessed through Solfy-generated feedback. No external measurement tools (e.g., acoustic analysis software) were employed, and results are presented as estimates. This mixed-method approach provided a comprehensive view of the effects generated by integ-

rating the application into the learning process.

The analysis of the collected data revealed a marked and consistent improvement in the participant's vocal performance across the intervention. During the initial stage, which relied exclusively on traditional methods, progress was modest: breath control improved slightly, and isolated sounds were intoned more accurately. However, persistent difficulties with tonal stability and fluent repertoire performance highlighted the limitations of training based solely on conventional techniques.

By contrast, the introduction of the Solfy application represented a turning point in the learning trajectory. Instant visual and auditory feedback enabled rapid error correction, while application-generated reports provided an objective basis for monitoring progress. As shown in Figure 1, the percentage of correctly intoned notes increased from an estimated 30% in Lesson 1 (Recording #35) to 100% in Lesson 12 (Recording #558), with incorrect and approximate notes eliminated. This progression demonstrates a substantial gain in vocal-technical accuracy. Percentages shown in Figure 1 are derived from Solfy's visual feedback reports and should be understood as estimates of pitch accuracy, rather than precise acoustic measurements.

The Solfy performance scores further corroborate these results. As illustrated in Figure 2, repeated attempts during early May (Recordings #241–388) consistently scored 0%, reflecting intonational instability. How-

Figure 1. Progression of Pitch Accuracy in Solfy.
Source: Authors

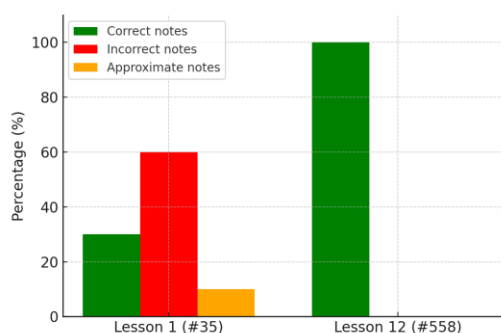
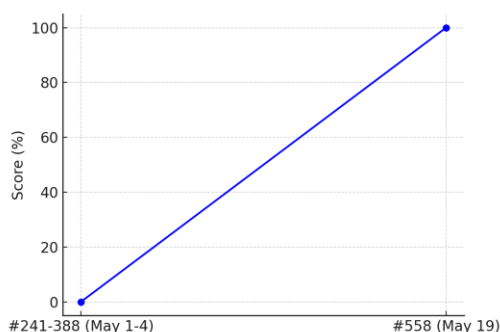


Figure 2. Progression of Solfy Performance Scores.
Source: Authors



ever, by May 19 (Recording #558), the participant achieved a score of 100%, indicating complete alignment with pitch targets. This sharp improvement suggests that iterative practice, reinforced by instant feedback, supported rapid consolidation of vocal control.

Beyond technical accuracy, the intervention fostered learner autonomy and self-regulation, directly addressing the evaluation of self-

monitoring and feedback impact. The possibility of repeating exercises until mastery, combined with visual progress tracking and orchestral rewards, created a strong sense of control over personal development. In her reflective journal, the participant described Solfy practice as “clearer, faster, and more motivating,” a finding consistent with recent studies emphasizing AI’s role in supporting self-regulation and reducing performance anxiety (Li et al., 2025; Zhang & Zhang, 2024). The results also confirm that AI complements rather than replaces the teacher’s role - so this method, corresponding to the needs of pedagogical transferability. The intervention demonstrated that traditional methods such as breathing and solfège remain essential but are amplified by instant feedback and self-monitoring. In this respect, Solfy acted as a “pedagogical mediator,” enabling independent practice without undermining corrective rigor. This conclusion echoes Pop-Sârb’s (2021) emphasis on the value of digital technologies as pedagogical support and resonates with Popean’s (2022) caution against over-reliance on technology at the expense of artistic authenticity.

Similar to recent studies and global trends, the present case shows that AI platforms can accelerate learning (Merchán Sánchez-Jara et al., 2024), stimulate autonomy (O’Leary, 2025), and democratize music literacy through accessible interactive resources (Muntean, Weidenfeld, & Koren, 2022). Nonetheless, the experiences also highlight contextual factors—such as the learner’s initial competence, selected repertoire, and teacher

readiness—that mediate effectiveness. These convergences underscore that findings from a Romanian pre-vocational context are consistent with global trends, suggesting transferability across cultural settings.

Moreover, the experiences support the view that AI tools like Solfy can function as vectors of music literacy. By combining singing with (self-)evaluation, the application extends beyond technical skill development to foster music reading and score comprehension. This aligns with ISME 2022 findings (Muntean & Koren), which demonstrated that self-singing solfège contributes to democratizing music literacy from the early years of schooling.

In summary, the experiences converge on two key directions:

- The responsible integration of AI in music education can accelerate technical progress and strengthen learner autonomy.
- The complementarity of traditional methods and digital applications provides a foundation for an integrative pedagogical model, where music literacy becomes more accessible, personalized, and reflective.

Conclusions

The AI-based Solfy application can produce positive effects on both vocal-technical progress and learner autonomy at the early stages of vocational training. Instant feedback, self-monitoring, and reward mechanisms contributed to an accelerated pace of learning and enhanced confidence in personal abilities. The experiences confirm that digital solutions do not replace traditional methods of music

education but rather complement them, amplifying the teacher's role. Instead of diminishing the importance of the instructor, applications such as Solfy extend the range of pedagogical intervention and create favorable conditions for differentiated and personalized instruction. This perspective aligns with international literature emphasizing the role of AI in accelerating learning, fostering self-regulation, and democratizing music literacy (Li et al., 2025; Merchán Sánchez-Jara et al., 2024; Muntean, Weidenfeld & Koren, 2022).

The introduced example demonstrated, music literacy can be advanced by combining singing with AI-assisted (self-)evaluation, thereby opening opportunities for broader democratization of access to music literacy. AI technologies can simultaneously support the development of musical and digital competences, an aspect particularly relevant within the current European and international educational landscape.

Several directions for future development emerge from these findings. Methodologically, larger-scale studies are needed to confirm the identified tendencies and to build scalable pedagogical models. Practically, AI integration could be explored beyond vocal training, extending to instrumental accompaniment, improvisation, and AI-assisted creativity. Organizationally, the study underscores the necessity of continuous professional development programs to prepare teachers for the responsible and effective use of digital tools. From the aspects of future research, the AI impacts on diagnostics and development in the area of music pedagogy,

and the systematic integration of AI into formal curricula can be investigated, thereby enriching classroom practice and advancing the dialogue on music education in the digital era.

Overall, the contribution of this study is to demonstrate that tradition and innovation are not mutually exclusive but can be articulated within an integrative pedagogical framework. The integration of AI applications such as Solfy offers a genuine opportunity to design music education that is more accessible, more personalized, and more reflective, attuned to the cultural realities and educational demands of the twenty-first century.

Authors' Note

This article was prepared by the authors, with support from AI-assisted tools used exclusively for language refinement, bibliographic verification, and editorial consistency. All conceptual design, analysis, and interpretation remain the authors' original contribution.

References

Solfy (website). Opened: 1.12.2025. URL:

<https://4solfy.com/>

Cheng, L. (2025). The impact of generative AI on school music education: Challenges and recommendations. *Arts Education Policy Review*. Doi: <https://doi.org/10.1080/10632913.2025.2451373>

European Commission. (2021). *Digital Education Action Plan 2021–2027: Resetting education and training for the digital age*. Publications Office of the European Union. Doi:

<https://doi.org/10.2766/149764>

Han, Y. (2025). Exploring a digital music teaching model integrated with recurrent neural networks under artificial intelligence. *Scientific Reports*, 15, 7495. Doi: <https://doi.org/10.1038/s41598-025-92327-8>

Hu, Z., Liu, Y., Zhang, Z., Chen, G., Yu, B. X. B., Li, J., & Cao, J. (2025). MusicScaffold: Bridging machine efficiency and human growth in adolescent creative education through generative AI (arXiv:2509.10327). arXiv. Opened: 1.12.2025. URL: <https://arxiv.org/abs/2509.10327>

Jin, L., Lin, B., Hong, M., Zhang, K., & So, H. J. (2025). Exploring the impact of an LLM-powered teachable agent on learning gains and cognitive load in music education (arXiv:2504.00636). arXiv. Opened: 1.12.2025. URL: <https://arxiv.org/abs/2504.00636>

Li, W., Cui, X., Manoharan, P., Dai, L., Liu, K., & Huang, L. (2025). AI-assisted feedback and reflection in vocal music training: Effects on metacognition and singing performance. *Frontiers in Psychology*, 16, 1598867. Doi: <https://doi.org/10.3389/fpsyg.2025.1598867>

Merchán Sánchez-Jara, J. F., González Gutiérrez, S., Cruz Rodríguez, J., & Syroyid, B. (2024). Artificial intelligence-

- assisted music education: A critical synthesis of challenges and opportunities. *Education Sciences*, 14(11), 1171. Doi: <https://doi.org/10.3390/educsci14111171>
- Muntean, L., & Koren, M. (2022). Self-Singing Solfege and (auto)evaluation – Music literacy for all? In I. Harvey & K. Rees (Eds.), *Proceedings of the 35th ISME World Conference on Music Education* (pp. 216–222). ISME.
- Muntean, L., Weidenfeld, S., & Koren, M. (2022). Promoting music literacy since primary school. In M. G. Popescu (Coord.), *Studii de muzicologie* (Vol. 17, pp. 76–89). Iași: Editura Artes. ISBN 978-606-547-455-0.
- Nichols, B. D. (2014). *The effect of SmartMusic on student practice* (Doctoral dissertation). Kennesaw State University.
- O’Leary, E. (2025). Considering the possibilities and problems of AI in music education: The need for critical literacies. *Action, Criticism, and Theory for Music Education*, 24(3), 138–164. Doi: <https://doi.org/10.22176/act24.3.138>
- Popean, M., & Căpățână, A. (2022). The Apocalypse of Authenticity: AI Faux Music / Apocalipsa autenticității. *Tehnologii Informaționale și de Comunicație în Domeniul Muzical*, 14(2), 83–97. Opened: 1.12.2025. URL: [https://tic.edituramediamusica.ro/reviste/2022/2/08%20Popean%20Capatana%20-%20ICTMF ISSN 2067-9408 2022 vol 14 issue 2.pdf](https://tic.edituramediamusica.ro/reviste/2022/2/08%20Popean%20Capatana%20-%20ICTMF%20ISSN%202067-9408%202022%20vol%2014%20issue%202.pdf)
- Pop-Sârb, D. (2021). Solfy—O abordare digitală pentru practica solfegiului în școală. *Tehnologii Informaționale și de Comunicație în Domeniul Muzical*, 12(1), 21–32. Opened: 1.12.2025. URL: [https://tic.edituramediamusica.ro/reviste/2021/1/ICTMF ISSN 2067-9408 2021 vol 12 issue 1 pg no 021-032.pdf](https://tic.edituramediamusica.ro/reviste/2021/1/ICTMF%20ISSN%202067-9408%202021%20vol%2012%20issue%201%20pg%20no%20021-032.pdf)
- UNESCO. (2023). *Guidance for generative AI in education and research*. Paris: UNESCO Publishing. Opened: 1.12.2025. URL: <https://unesdoc.unesco.org/ark:/48223/pf0000386047>
- Wei, X., Zhang, Y., & Li, J. (2022). College music education and teaching based on AI techniques. *Computers & Electrical Engineering*, 100, 107851. Doi: <https://doi.org/10.1016/j.compeleceng.2022.107851>
- Zhang, F., & Zhang, P. (2024). Transforming music education through artificial intelligence: A systematic literature review. *International Journal of Interactive Mobile Technologies (ijIM)*, 18(18), 76–93. Doi: <https://doi.org/10.3991/ijim.v18i18.50545>

MŰHELY, RENDEZVÉNY
WORKSHOPS AND EVENTS

MI ÉS AZ NTP-TEHETSÉG-25-0142 PROJEKT

Author(s) / Szerző(k):

Szűts Zoltán (Prof., Ph.D.)
Eszterházy Károly Katolikus Egyetem
(Magyarország)

Szőke-Milinte Enikő (Ph.D.)
Eszterházy Károly Katolikus Egyetem
(Magyarország)

E-mail:

szuts.zoltan@uni-eszterhazy.hu

Cite: Idézés:

Szűts Zoltán és Szőke-Milinte Enikő (2025): MI és az NTP-TEHETSÉG-25-0142 projekt. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, VII. évf. 2025/2. szám. 129-133.
Doi: <https://www.doi.org/10.35406/MI.2025.2.129>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE:

Ethics Permission / Etikai engedély: KFS/2025/MI0016

Reviewers: Lektorok:

Public Reviewers / Nyilvános Lektorok:
1. Szűts Zoltán (Prof., Ph.D.), Eszterházy Károly Katolikus Egyetem
2. Mogyorósi Zsolt (Ph.D.), Eszterházy Károly Katolikus Egyetem

Anonymous reviewers / Anonim lektorok:
3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

A Nemzeti Tehetségprogram és a Kulturális és Innovációs Minisztérium összesen 3 590 000 forintos támogatásával az Eszterházy Károly Katolikus Egyetem (EKKE) Pedagógiai Kara a 2025/2026-os tanévben valósítja meg a „Pedagógia a mesterséges intelligencia fényében és árnyékában (nem csak) a STEM tantárgyak esetében” című tehetséggondozó programot (pályázati azonosító: NTP-TEHETSÉG-25-0142). Ez a tanulmány a projekt mesterséges intelligencia (MI) technológiával való kapcsolatát foglalja össze.

Kulcsszavak: mesterséges intelligencia, tehetség

Diszciplínák: pedagógia, pszichológia, informatika

Abstract

AI AND THE NTP-TEHETSÉG-25-0142 PROJECT

With the support of the National Talent Program and the Ministry of Culture and Innovation totaling HUF 3,590,000, the “Pedagogy in the Light and Shadow of Artificial Intelligence (not only) in the Case of STEM Subjects” talent management program (application ID: NTP-TEHETSÉG-25-0142) will be implemented in the 2025/2026 academic year by the Faculty of Pedagogy of the Eszterházy Károly Catholic University (EKKE). This study summarizes the connection of the project with artificial intelligence (AI) technology.

Keywords: artificial intelligence, talent

Disciplines: pedagogy, psychology, informatic

Az NTP-TEHETSÉG-25-0142 pályázati azonosítószerű projekt címe: „Pedagógia a mesterséges intelligencia fényében és árnyékában (nem csak) STEM tárgyak esetében”. A projekt az [Eszterházy Károly Katolikus Egyetem Pedagógiai Kara](#) szervezésében valósul meg 2025. szeptember 1. és 2026. augusztus 31. között a [Nemzeti Tehetség Program](#) és a [Kulturális és Innovációs Minisztérium](#) 3 590 000 Ft összegű támogatásával. Az alábbiakban a 60

órás complex tehetséggondozó programot magába foglaló projekt mesterséges Intelligencia (MI) technológiával kapcsolatos vonatkozásai kerülnek összefoglalásra.

Digitális kompetenciafejlesztés – különös tekintettel a mesterséges intelligencia használatára

A projekt révén megvalósuló dúsító, gazdagító tehetséggondozó program egyik

lényeges célja, hogy a résztvevők felkészítése arra, hogy mire és miként tudják hatékonyan és etikusán használni a napjainkban ingyenesen elérhető mesterséges intelligencia alapú alkalmazásokat: a) középiskolás tanulóként, b) leendő (talán pedagógusképzésben résztvevő) egyetemi hallgatóként, c) pedagógusként (amennyiben valóban pedagóguspályára kerülnek).

Cél, hogy a résztvevők megismerjék a mesterséges intelligencia használatával járó előnyöket, hátrányokat, veszélyeket és lehetőségeket – annak érdekében, hogy a mesterséges intelligencia etikus és szakszerű felhasználóivá válhassanak akár tanulói/hallgatói, akár pedagógus/oktató szerepkörökben is.

Módszertan

A mesterséges intelligencia (MI, artificial intelligence) a tanulók, hallgatók, pedagógusok számára egyaránt hasznos eszköznek tekinthető, melynek szakszerű és etikus használatára azonban még képezni szükséges a felsorolt potenciális felhasználókat. Az Eszterházy Károly Katolikus Egyetem Pedagógiai Kara oktatóinak mesterséges intelligencia tárgyú publikációi, kutatási és oktatási tapasztalatai (lásd például Prof. dr. Szűts Zoltán, Dr. Mező Ferenc publikációs listáját a www.mtmt.hu oldalon) jelen pályázat keretében is hasznosításra kerülnek.

A tehetséggondozó foglalkozás-sorozatban digitális kompetenciafejlesztés egy-

részt a szövegszerkesztés, táblázatkezelés és prezentációkészítést segítő szoftverek használatára fókuszál, másrészt a virtuális kiállítások tanulási/tanítási használati lehetőségeit járja körül, harmadrészt a mesterséges intelligenciával kapcsolatos alábbi témaköröket érinti:

- 1) A mesterséges intelligencia alapfogalmai;
- 2) Mesterséges intelligencia a STEM tárgyakkal (is) kapcsolatos forráskutatásban;
- 3) Mesterséges intelligencia a STEM tárgyakkal (is) kapcsolatos kivonatok készítésében;
- 4) Mesterséges intelligencia a STEM tárgyakkal (is) kapcsolatos ismeretek(ön)ellenőrzésében;
- 5) Mesterséges intelligencia a STEM tárgyakkal (is) kapcsolatos prezentációk összeállításában.

A minimálisan szükséges elméleti alapokat előadás formájában kapják meg a résztvevők, de a digitális kompetenciafejlesztés döntő részben (90%-ban) gyakorlati jellegű, a szoftverek, mesterséges intelligencia alkalmazások közvetlen használatán alapul. E kompetenciákat összetett módon próbálhatják ki a gyakorlatban, például a projekt keretében számukra biztosított nemzetközi interdiszciplináris konferenciákon. Így például 2025. december 5.-én lehetőség nyílt arra, hogy 30 tanuló részt vehessen a „Kreativitás – Elmélet és gyakorlat (2025)” Nemzetközi Interdiszciplináris Konferencián (Mező, 2025), a-

hol nemcsak a tudományos stílusú eszmecserét, hanem az azzal kapcsolatos digitális eszközhasználatot és a mesterséges intelligencia alkalmazást is megfigyelhették.

A digitális kompetencia fejlesztése (mesterséges intelligencia tanulás/tanítási célú használata, tanulás-, tanítás- és kutatómódszertani fejlesztés a STEM tárgyak vonatkozásában, szövegszerkesztő, táblázatkezelő, prezentációkészítő, virtuális kiállítás-kezelő szoftverek tekintetében) tehát sajátélmény biztosításával történik meg.

A foglalkozássorozat nem titkolt célja a mesterséges intelligenciával kapcsolatos pozitív irányú attitűdformálás és kompetenciafejlesztés is.

Záró megjegyzések

A mesterséges intelligenciára fókuszáló programelemek hatására a résztvevők etikusan és hatékonyan fognak tudni használni:

- a) szövegszerkesztő szoftvert;
- b) táblázatkezelő szoftvert;
- c) prezentációkészítést segítő szoftvert;
- d) virtuális kiállítást kezelő szoftvert;
- e) mesterséges intelligencia alapú alkalmazásokat.

E digitális jártasságok és mesterséges intelligenciát érintő felhasználói tapasztalatok mind a STEM, mind a pedagó-

gusképzés felé orientálódó tanulók számára jól hasznosíthatók a jövőben.

Lényeges továbbá, hogy a projektben résztvevő tanulók énképébe épüljön be, hogy ők a pedagógiai pálya, azzal kapcsolatban a STEM tárgyak és a mesterséges intelligencia használata iránt érdeklődő, e területeken legalább kortársaikhoz képest kimagasló teljesítményeket elérni képes személyek, akiket e tulajdonságaik, teljesítményeik miatt (is) szociális környezetük (kortársak, pedagógusok, család) tehetségesnek tart.

Köszönetnyilvánítás

A „Pedagógia a mesterséges intelligencia fényében és árnyékában (nem csak) STEM tárgyak esetében” című (NTP-TEHETSÉG-25-0142 azonosítószámú) projektet a [Nemzeti Tehetség Program](#) és a [Kulturális és Innovációs Minisztérium](#) 3 590 000 Ft összeggel támogatta. A támogatást ezúton is köszönjük!



KULTURÁLIS ÉS INNOVÁCIÓS
MINISZTERIUM

Irodalom

Eszterházy Károly Katolikus Egyetem

Pedagógiai Kar honlapja. Megnyitva:

2025.12.10. URL: [https://uni-](https://uni-eszterhazy.hu/pk)

[eszterhazy.hu/pk](https://uni-eszterhazy.hu/pk)

Kreativitás – Elmélet és gyakorlat (2025)

Nemzetközi Interdiszciplináris Konferencia

(hivatalos weboldal). Megnyitva:

2025.12.05. URL:

<https://www.kpluszf.com/crea-conf-2025/>

Kulturális és Innovációs Minisztérium

(hivatalos honlap). megnyitva:

2025.12.10. URL:

<https://kormany.hu/kormanyzat/kulturalis-es-innovacios-miniszterium>

Mező Ferenc (2025): *Program – Kreativitás*

– *Elmélet és gyakorlat (2025) konferencia.*

Megnyitva: 2025.12.05. URL:

https://drive.google.com/file/d/1jew6AAAsYF8Jr26iMY4tYdHx_NnbC0U21/view

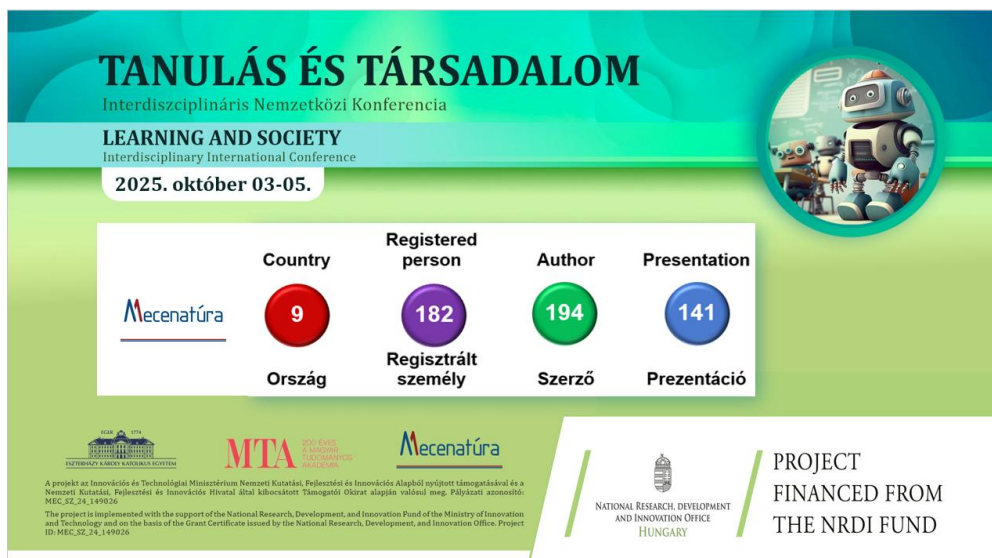
Nemzeti Tehetség Program (hivatalos

honlap). Megnyitva: 2025.12.10. URL:

<https://nemzetitehsegprogram.hu/>

SHORT REPORT ABOUT THE ‘LEARNING AND SOCIETY (2025)’ CONFERENCE AND THE MEC-SZ-24-149026 PROJECT

Mező, Ferenc (Ph.D.)



“Learning and Society (2025)” International Interdisciplinary Conference was held in Eger, between October 3-5, 2025, in the main building of Eszterházy Károly Catholic University, as one of the events aimed at celebrating the 200th anniversary of the Hungarian Academy of Sciences.

The project is implemented with the support of the National Research, Development, and Innovation Fund of the Ministry of Innovation and Technology and on the basis of the Grant Certificate

issued by the National Research, Development, and Innovation Office. Project ID: MEC-SZ-24-149026.

About the conference

Based on preliminary registration data, 182 people from at least nine countries (Cambodia, the Netherlands, China, Hungary, Malaysia, Romania, Serbia, Slovakia and Ukraine) registered for the event, at which 194 authors (including 12 who were not registered by the conference

but were listed as co-authors) presented a total of 141 presentations in 13 sections. The presentations were public.

Four public plenary lectures were:

- Horák, Rita (Prof. PhD; Serbia): The Responsibility of Education in Environmental Protection: Perspectives from Serbia.

- Nagy, Melinda (Dr. Habil. PaedDr.; Slovakia): Learning and Society: Expert Dialogues and Community Engagement in Addressing Local Challenges

- Pásztor, Rita (PhD; Romania) and Gál, Katalin (PhD; Romania): Promoting Student Success and Equal Opportunities: Needs Assessment for the Foundation of Disadvantage Compensation Programs in Secondary and Higher Education

- Mező, Ferenc (PhD; Hungary) and Mező, Katalin (PhD; Hungary): The OxIPO Project: increasing the human information processing performance

Accompanying programs were:

- 2025.10.03-05.: Art and Poster Exhibition

- 2025.10.03., 18.10-18.30: Military history model exhibition

- 2025.10.04., 9.00-10.00: Filming, interviews

- 2025.10.04., 16.30-17.30: Science show in the astronomy tower

- 2025.10.04., 18.00-19.00: International Learning Research Network

- 2025.10.04., 20.00-22.00: Conversation circle with Dudás Éva, Halasi Szabolcs, Juhász Dániel, Mező Ferenc, Nagy Kanász

Viola, Papp Zoltán, Pintér Krekic Valéria, Udvaros Anna, Udvaros Borbála, Varga Zsolt.

- 2025.10.05., 10.30-12.00: “Innovation – from the aspect of learning and society” workshop

- 2025.10.05., 14.00-16.00: Organizing Committee Meeting

Organizing Committee

The Head of the Organizing Committee was Ferenc Mező (Ph.D., Eszterházy Károly Catholic University, Hungary).

The Members of the Scientific Committee were:

- Gál, Katalin (Ph.D., Partium Christian University, Romania)

- Hanák, Zsuzsanna (Habil., Ph.D., Eszterházy Károly Catholic University, Hungary)

- Horák, Rita (Prof., Ph.D., University of Novi Sad, Serbia)

- Mező, Katalin (Ph.D., University of Debrecen, Hungary)

- Mogyorósi, Zsolt (Ph.D., Eszterházy Károly Catholic University, Hungary)

- Muntean, Loredana (Ph.D., Universitatea din Oradea, Romania)

- Pšenáková, Ildikó (Habil., Ph.D., Trnava University in Trnava, Slovakia)

- Szabó, Tibor (Ph.D., Constantine The Philosopher University in Nitra, Slovakia)

- Szabóné, Balogh Ágota Márta (Ph.D., Gál Ferenc University, Hungary)

- Szűts, Zoltán (Ph.D., Eszterházy Károly Catholic University, Hungary)
- Ujfaludi, László (Ph.D., Prof.emeritus, Eszterházy Károly Catholic University, Hungary)
- Varga, Zsolt (Ph.D., University of Debrecen, Hungary)

The members of the event organizing group were the next coworkers of the Eszterházy Károly Catholic University (Hungary)

- Gál, Tibor (Eszterházy Károly Catholic University, Hungary)
- Kormos, Ádám (Eszterházy Károly Catholic University, Hungary)
- Oszlanczi, Krisztina (Eszterházy Károly Catholic University, Hungary)
- Tóth, Ádám (Project Manager, Project Management Group)
- Román, Enikő (Finance Manager, Project Finance Group)
- Varga, Ferenc (Head of Department, Communication Department)
- Berecz, Adrienn (Event Planner, Communications Group)
- Bíró, Tünde (Graphic designer)
- Izsóf, Rebeka (Event Planner, Communications Group)

The volunteer student helpers were the following psychology students of the Eszterházy Károly Catholic University (Hungary):

- Csajbók, Gábor
- Hegedűs, Marcell

- Szabó, Klára
- Cz. Tóth, Csenge Flóra
- Bodnár, Alíz
- Szentgyörgyi, Kitti

Chairmen of the sections were:

- Section 1: Szabóné Balogh, Ágota (PhD) and Szabó, Tibor (PhD)
- Section 2: Szűts, Zoltán (Prof., PhD) and Mogyorósi, Zsolt (PhD)
- Section 3: Gál, Katalin (PhD) and Hanák, Zsuzsanna (PhD, Habil.)
- Section 4: Horák, Rita (Prof., PhD) and Mező, Katalin (PhD)
- Section 5: Müller, Anetta (Prof., PhD) and Urbinné Borbély, Szilvia (PhD)
- Section 6: Mező, Katalin (PhD) and Mező, Ferenc (PhD)
- Section 7: Bogárdi, Tünde (PhD) and Frank, Tamás (PhD)
- Section 8: Pšenáková, Ildikó (PhD, Habil.) and Simándi, Szilvia (PhD, habil.)
- Section 9: Pšenáková, Ildikó (PhD, Habil.) and Pšenák, Peter (PhD)
- Section 10: Vida, Gergő (PhD) and Mikelayi, Wumaier (Drs)
- Section 11: Ujfaludi, László (Prof. CSc.) and Muntean, Loredana (PhD)
- Section 12: Lestyán, Erzsébet (PhD) és Kós, Nóra (PhD)
- Section 13: Mező, Ferenc (PhD)

Products

Products of the MEC_SZ_24_149026 project are:

- 1) The URL of the project website is:

<https://uni-eszterhazy.hu/tanulas-konferencia/hirek>

2) "Learning and Society" Interdisciplinary International Conference (Eger, October 3-5, 2025). 194 authors from 9 countries took part in the event, and 141 presentations were presented in 13 sections.

3) The e-brochure containing the program and abstracts of the "Learning and Society" Interdisciplinary International Conference was published. Bibliographic data: Mező Ferenc (Szerk.) (2025): *Tanulás és társadalom (2025) – Program és absztraktok*. EKKE, Eger. URL: https://publikacio.uni-eszterhazy.hu/8888/1/absztr_teljes.pdf

4) A 152-page (A4 size) e-publication summarizing the project was published. Bibliographic data: Mező Ferenc (Szerk.) (2025): *Tanulás és társadalom (2025) – Válogatott tanulmányok*. EKKE, Eger. URL: <https://publikacio.uni-eszterhazy.hu/8883/1/tanulmanykotet.pdf>

5) The short film related to the Interdisciplinary International Conference "Learning and Society" was released on the project website. URL: <https://www.youtube.com/watch?v=sdnZc92kmgw>

6) 14 online news reports about the project have been published on the project website.

7) Invitations to the "Learning and Society (2025)" conference were published in the following journals and issues:

- Különleges Bánásmód folyóirat, 2024/4.

- Különleges Bánásmód folyóirat, 2025/1.
- Különleges Bánásmód folyóirat, 2025/2.
- Különleges Bánásmód folyóirat, 2025/3.
- Különleges Bánásmód folyóirat, 2025/SI
- OxIPO folyóirat, 2024/4.
- OxIPO folyóirat, 2025/1.
- OxIPO folyóirat, 2025/2.
- OxIPO folyóirat, 2025/3.
- Mesterséges intelligencia folyóirat, 2024/2.
- Mesterséges intelligencia folyóirat, 2025/1.
- Lélektan és hadviselés folyóirat, 2024/2.
- Lélektan és hadviselés folyóirat, 2025/1.

8) A report on the event was published in the following journals:

- OxIPO folyóirat, 2025/4.
- Mesterséges intelligencia folyóirat, 2025/2.
- Lélektan és hadviselés folyóirat, 2025/2.

Acknowledgements

The project is implemented with the support of the National Research, Development, and Innovation Fund of the Ministry of Innovation and Technology and on the basis of the Grant Certificate issued by the National Research, Development, and Innovation Office. Project ID: MEC-SZ-24-149026



A GÁBOR DÉNES EGYETEM MIT ÉS A STRT HOLDING INGYENES ONLINE MI TANANYAGA

Forrás: <https://mijovunk.hu/>



Gábor Dénes Egyetem
Mesterséges Intelligencia Tudásközpont

A mesterséges intelligencia egyre nagyobb hatással van az életünkre, a munkánkra és a közigazgatásra is. Hamarosan a mindennapjaink szerves része lesz – ezért fontos, hogy Te is felfedezd, milyen lehetőségeket rejt számodra.

Készítettünk egy ingyenes kurzust, amellyel mindössze 3 óra alatt megismerheted, megértheted és elkezdheted hatékonyan használni ezt a forradalmi technológiát. A végén még mikrotanúsítványt is szerezhetsz.

A 3 órás ingyenes képzés 2025.06.06. – 2026.08.31. között érhető el a „Regisztráció” menüponton keresztül:

<https://mijovunk.hu/>

A nyitva álló idő alatt a képzés bármikor elvégezhető.

A képzés a felnőttképzésről szóló 2013. évi LXXVII. (továbbiakban: Fktv.) törvény szerinti felnőttképzés, mely mikrotanúsítvánnyal zárul.

A felnőttképzési mikrotanúsítvány olyan rövid időtartamú képzés elvégzését igazoló dokumentum, amely képzésnek a tananyaga a szakképzési törvény szerinti tananyagjegyzékben szerepel.

Ezeknek a képzéseknek a munkaerőpiaci értékét nem a konkrét képző intézmény, hanem a mikrotanúsítványt adó képzés akkreditált tananyaga adja, ugyanakkor a mikrotanúsítvány szakképesítést és szakértettséget nem igazol.

Miről fogok tanulni?

 Bevezető 45 perc Hogyan jutottunk idáig? Az MI története, működése és jelenlegi helyzete.	 MI eszközök 90 perc Hogyan használd az MI-t munkában, tanulásban és a hétköznapiakban – eszközök, promptolás, gyakorlati példák.	 Az MI hatása 45 perc Milyen hatása van az MI-nek a világra? Étika, biztonság, gazdaság és a jövő kihívásai.
---	--	--

Bejelentkezés

A képzés iránt érdeklődő személy regisztrál a www.mijovunk.hu weboldalon az ott meghatározott feltételek mentén.

Regisztrációs folyamat

A képzés iránt érdeklődő személy a www.mijovunk.hu weboldalon elfogadja az ÁSZF-et, megadja a személyes adatait, valamint megteszi a szükséges nyilatkozatokat az Fktv.-ben és az ÁSZF-ben leírtak szerint. A regisztrációs folyamat végén a felnőttképző a képzésben részt vevő által megadott e-mail címre megküldi a képzés elektronikus elérhetőségét, azonban fontos információ, hogy ez még nem minősül a jelentkezési folyamat eredményes véglegesítésének.

A jelentkezési folyamat eredményes véglegesítése: A jelentkezés során megadott adatok az Fktv. szerinti felnőttképzési adatszolgáltatás keretében ellenőrzésre kerülnek, és amennyiben ez az elektro-

nikus ellenőrzés sikeres volt, és az adatok megfelelnek a valóságnak, úgy a sikeres jelentkezésről a rendszer visszaigazolást küld a jelentkezőnek.

Képzés ismertető magánszemélyeknek

Bizonytalan vagy elutasító vagy az MI-vel kapcsolatban? Netán felismerted már a benne rejlő lehetőségeket, és szeretnéd másokkal is megosztani ezt a tudást, meggyőzni környezetet az MI megismerésének fontosságáról?

A Gábor Dénes Egyetem (GDE) mesterséges intelligencia ingyenes online képzése segítséget nyújt Neked, hogy felfedezd a Mesterséges Intelligenciát (MI), vagy a már meglévő tudásod birtokában másokkal is hitelesen megismertessd ezt a rendkívül izgalmas új világot.

Intézményünk (GDE), az online képzések szakértője, most egy különleges lehető-

séget kínál: az egyetemünk Nexius rendszeréből bárholnan elérhető, háromórás online tananyag bevezet a mesterséges intelligencia különleges és rohamosan fejlődő világába. Az ingyenes képzés elvégzésével felnőttképzési mikrotanúsítványt szerezhetsz, amely igazolja az MI képzésünk elvégzését és a tudás elsajátítását.

Mit kínál a tananyag azoknak, akik elvégzik a képzést?

- Első rész – Az MI alapjai és fejlődése: Megismerheted, mi is valójában a mesterséges intelligencia, és hogyan jutottunk el a mai, lenyűgöző eredményekig.

- Második rész – A promptolás művészete: Felfedezheted a promptolás (az MI-nek adott utasítások) logikáját és egyszerű alkalmazhatóságát, továbbá számos gyakorlati alkalmazáson keresztül – többek között az oktatásban, a munkahelyen és a mindennapokban – képet kapsz ennek hasznosságáról.

- Harmadik rész – Az MI a gyakorlatban és a jövőben: Betekintést nyersz abba, hogyan használják már ma az MI-t olyan területeken, mint az ügyvédi munka, az informatika, a szoftverfejlesztés, a bank-szektor vagy az orvostudomány, és milyen további fejlődésre számíthatunk. Ebben a részben hívjuk fel figyelmedet az MI-vel kapcsolatos esetleges veszélyekre és azok elkerülésének lehetőségeire, illetve a követendő etikai kérdésekre is.

Kinek szól ez a képzés?

Azoknak, akik:

- még nem ismerik az MI-t: A kurzus tökéletes kiindulópont, ha szeretnéd megérteni, miről szól a mesterséges intelligencia, és hogyan befolyásolja az életünket.
- elutasítják az MI-t: Adj mégis egy esélyt a megismerésnek! A tananyag segít eloszlatni a tévhitet, bemutatja az MI előnyeit és a kockázatok kezelését is.
- már ismerik az MI-t, és meggyőznék a környezetüket: A képzésből ötleteket meríthetsz, hogyan kommunikálj hatékonyan az MI-ről, és hogyan ösztönözz másokat is a tanulásra és az MI felősségteljes használatára.

Miért különösen hatékony a GDE ingyenes online MI képzése?

Mert a nemzetközi szinten a legkorszerűbb és legrugalmasabb tanulási platformok egyike, amelyben:

- az online tanfolyam és tananyag lehetővé teszi, hogy a résztvevők saját tempójukban és időbeosztásuk szerint tanuljanak, így a tanulás könnyen beilleszthető a mindennapi életbe – akár munka vagy család mellett is.
- a GDE Nexius keretrendszere biztosítja a kiváló minőségű videós előadásokat (oktató, prezentáció, és ha kell, akkor a tábla előadással szinkron képét is) interaktív funkciókkal (pl. többek között egyéni jegyzetelés, keresési lehetőség akár az oktató által kimondott szavakra is).

- az egyedülálló vizualitás, letölthető példák és illusztrációk segítik a megértést még azok számára is, akik kevésbé technikai beállítottságúak.

Az MI nem a gépek fenyegetése – hanem egy új eszköz a mindennapi élet, munka és tanulás támogatására.

A tananyag készítői

- Gábor Dénes Egyetem Mesterséges Intelligencia Tudásközpont: a több mint 30 éve az oktatás élvonalában levő - Gábor Dénes Egyetem hozta létre többek között azzal a céllal, hogy aktív, vezető szerepet vállaljon a mesterséges intelligenciával kapcsolatos tudásmegosztásban és edukációban.

- STRT Holding Nyrt.: Magyarország legaktívabb startup befektetője és meghatározó vezetőképzője, akik abban hisznek, hogy „egy ország sorsát hosszú távon a vállalkozói, vállalatai határozzák meg”.

**Ne halogasd,
csatlakozz most
a képzéshez, és
szerezzél
versenyelőnyt a jövő
digitális világába!**

FELHÍVÁS
INTERDISZCIPLINÁRIS JUNIOR KUTATÓCSOPORTBA
TÖRTÉNŐ BEKAPCSOLÓDÁSRA



Cél:

Középiskolások, BA, BSC, MA, MSC, PHD hallgatók számára lehetőséget adni a saját diszciplínájukon átívelő kutatásokba bekapcsolódásra, publikációkat megjelentetésére, nemzetközi konferenciárészvételre.

A bekapcsolódással járó haszon

A részvétel a bekapcsolódók számára azért hasznos, mert:

- a) ösztöndíjak, pályázatok során érvényesíthető teljesítményei (publikáció, konferenciaelőadás) lesznek,
- b) saját témájában kutathat és azt gazdagíthatja kutatótársai szaktudását is felhasználva,
- c) életrajzában is jól mutató bejegyzést kap,
- d) szakmai kapcsolatrendszere bővül,

- e) ingyen vehet részt nemzetközi konferenciákon,
- f) ingyen publikálhat Open Access (nyílt hozzáférésű) kiadványokban.

Feladatok

A résztvevő feladata a következő lesz:

- 1) Jelentkezés a csoportba (a felhívás végén látható linken keresztül)
- 2) A csoport alakuló ülésén (személyes vagy online) részvétel a közös kutatási téma kialakításában. Például: korábbi hasonló csoportban pszichológia, jogtudo-

mány, gazdaságtudomány és orvostudomány szakos hallgatók fordultak saját szakjuk felől közös érdeklődésbe vágó kérdésekhez.

3) 10 perces prezentációval ingyenes részvétel a minden év decemberében megrendezésre kerülő „Kreativitás – Elmélet és Gyakorlat Nemzetközi Interdiszciplináris Konferencia” című rendezvényen. Magyar vagy angol nyelvű előadásokat lehet majd tartani, amiről két nyelvű igazolást állítanak ki a Szervezők. Az előadások témáját Ön választhatja meg.

4) Min. 1 tanulmány megírása. A megjelentetést megegyezés szerint folyóiratban vagy szöveggyűjteményben tervezzük.

Kiket várunk a programba?

A jelentkezést azoknak a középiskolásoknak, hallgatóknak, doktoranduszoknak ajánljuk, akik:

- a) sokoldalúak, s kíváncsiak arra, hogyan tudnak együttműködni különböző tudományágak képviselőivel;
- b) teljesítmény-centrikusak: a részvétel publikációkkal, konferenciákon történő előadásokkal is jár;
- c) tudományos karrierjüket, s széleskörű kapcsolatrendszerüket már hallgatóként igyekeznek megalapozni;
- d) a hétköznapi hallgatói létet kellemes és hasznos időtöltéssel igyekeznek kiegészíteni;
- e) kedvelik a jó társaságot.

Részvételi díj

A programban való részvétel díj: 0 Ft.

A program keretében megrendezésre kerülő nemzetközi online konferenciákon történő részvételi díj: 0 Ft.

A programban történő folyóiratokban, tanulmánykötetben történő tanulmány megjelentetésének díja: 0 Ft.

A program egyéb költséget nem tartalmaz, de a résztvevők a saját kutatási munkájukkal kapcsolatban esetlegesen felmerülő költségeket önmállóan fedezik.

Időigény

A program időigénye: kb. 2 óra/alakuló megbeszélés + saját ütemű kutatás és publikáció írás + konferenciákon való részvétel.

Amit lehet, elektronikusan oldunk meg, ezzel csökkentve az időigényt.

Jelentkezési határidő:

2026. március 11.

Jelentkezés módja: bejelentkező e-mail küldése erre az e-mail címre:

ferenc.mezo1@gmail.com

Szervező és támogatók



E tehetséggondozó program a Kocka Kör Tehetséggondozó Kulturális Egyesület és a Nemzetközi Tanuláskutató Hálózat (ILEARN) együttműködése keretében valósul meg.

Kapcsolat, további információ

Szakmai vezető: Dr. Mező Ferenc

E-mail: ferenc.mezo1@gmail.com

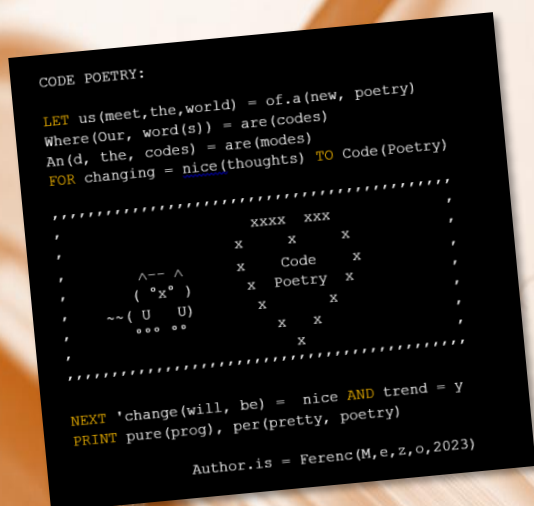
Mobil: 06 30 656 1 565

CODE POETRY PÁLYÁZAT (2026)

ÍRJ KÓDVERSET!

BÁRMILYEN PROGRAMNYELVET HASZNÁLHATSZ!

BÁRMIRŐL SZÓLHAT A VERS
(AMI NEM KIREKESZTŐ, NEM JOGSÉRTŐ).



A VERS TERJEDELME: MIN. 2 KÓDSOR

A KÓDVERSEKET 2026. SZEPTEMBER 31-IG KÜLDD EL AZ

INFO@KPLUSZF.COM

CÍMRE EGY RÖVID KÍSÉRŐ ÜZENETTEL, AMI TARTALMAZZA:

1. A SZERZŐ NEVÉT
2. A KÓDVERS CÍMÉT
3. A KÓDVERS PROGRAMNYELVÉT

A közlésre alkalmas kódverseket a K+F Stúdió Kft. (www.kpluszf.com) által kiadott e-kiadványban és/vagy e-folyóiratszámokban tesszük közzé, illetve angol-magyar kétnyelvű igazolást adunk a műről.

További információ az info@kpluszf.com e-mail címen keresztül kérhető.

A kódköltészetéről (code poetry), illetve a kódversekről (code poems) háttéranyag, módszertani útmutató található ezekben a cikkekben:

Mező Ferenc (2023): Code Poetry – avagy: Amikor az irodalom csókot dob az informatikának, de a mesterséges intelligencia elkapja azt a tehetséggondozás örömére... *Mesterséges intelligencia – interdiszciplináris folyóirat*, V. évf. 2023/1. szám. 9-19. doi: [10.35406/MI.2023.1.9](https://doi.org/10.35406/MI.2023.1.9)

Mező Ferenc (2023): Code Poetry – Módszertani javaslatok tehetségfejlesztő programok számára. *Mesterséges Intelligencia – interdiszciplináris folyóirat*, V. évf. 2023/1. szám. 87-98. doi: [10.35406/MI.2023.1.87](https://doi.org/10.35406/MI.2023.1.87)

A PÁLYÁZATRA TÖRTÉNŐ KÓDVERSEK BEKÜLDŐI A MŰ BEKÜLDÉSÉVEL NYILATKOZNAK ARRÓL, HOGY A KÓDVER S A SAJÁT SZELLEMI TERMÉKÜK, S HOZZÁJÁRULNAK ANNAK KÖZLÉSÉHEZ A K+F STÚDIÓ KFT. ÁLTAL KIADOTT E-KIADVÁNYOKBAN, E_FOLYÓIRATOKBAN.

**> DO_NOT_FORGET:
> PLEASE.WRITE CODE(POEMS)**