

**MACHINE E-LEARNING
IN ARTIFICIAL INTELLIGENCE –
COMPARATIVE STUDY****GÉPI TANULÁS A
MESTERSÉGES INTELLIGENCIÁBAN –
ÖSSZEHASONLÍTÓ TANULMÁNY**OLÁH Róbert¹**Abstract**

The study provides an insight into the field of machine learning (ML) by reviewing the literature of the past decades, primarily in Hungarian, and providing a detailed comparison between different theoretical approaches and applications. Based on the review of the source works, the study demonstrates that machine learning is not just a technological tool, but a multidisciplinary phenomenon that fundamentally shapes scientific and social structures. The article pays special attention to the three main methodological directions of machine learning, their practical applications, and the ethical and social challenges. The authors describe in detail the architecture of neural networks, the learning model, pathfinding on graphs, and supervised learning, where the quality of annotation affects the accuracy of algorithms. The study compares the work of several Hungarian authors along the formal frameworks, the applied methods, and the theoretical differences, highlighting the common foci of the research.

Keywords

Machine learning, Supervised learning, Unsupervised learning, Reinforcement learning, Artificial neural networks, Deep learning

Absztrakt

Ez a tanulmány a gépi tanulás (ML) területére nyújt betekintést az elmúlt évtizedek elsősorban magyar nyelvű szakirodalmak feldolgozásával, részletes összehasonlítást nyújtva a különböző elméleti megközelítések és alkalmazások között. A forrásmunkák áttekintése alapján a tanulmány bemutatja, hogy a gépi tanulás nem csupán egy technológiai eszköz, hanem egy multidiszciplináris jelenség, amely alapvetően alakítja a tudományos és társadalmi struktúrákat. A cikk különös figyelmet fordít a gépi tanulás három fő módszertani irányára, azok gyakorlati alkalmazásaira, valamint az etikai és társadalmi kihívásokra. A szerzők részletesen ismertetik a neurális hálózatok architektúráját, a tanulási modellt, a gráfokon történő útkeresést és a felügyelt tanulást, ahol az annotáció minősége befolyásolja az algoritmusok pontosságát. A tanulmány több hazai szerző munkásságát veti össze a formai keretek, az alkalmazott módszerek és az elméleti különbségek mentén, rávilágítva a kutatások közös fókuszpontjaira.

Kulcsszavak

Gépi tanulás, Felügyelt tanulás, Felügyelet nélküli tanulás, Megerősítéssel tanulás, Mesterséges neurális hálózatok, Mélytanulás

¹ olah.robert.oe@gmail.com | ORCID: 0009-0007-9134-9521 | IT teacher, Educational Center of Érd | Informatikai oktató, Érdi Tankerületi Központ

GÉPI TANULÁS KATEGÓRIÁK

Gépi tanulás kategóriák

A vizsgált forrásdokumentumok alapján a gépi tanulás (ML) nem csupán egy technológiai eszköz, hanem a mesterséges intelligencia egyik legdinamikusabb ága, amely napjainkban számos tudományterületet formál át és nyer alkalmazást. A szakirodalmi kutatás alapján megállapítható, hogy a szerzők egységesek a gépi tanulás három fő kategóriába sorolásában:

- **Felügyelt tanulás (Supervised Learning):** Címkeztet adatokból (példapárokból) tanul, célja a klasszifikáció (osztályzás) vagy regresszió. A szociológiában például gyűlöletbeszéd felismerésére, a marketingben lemorzsolódás előrejelzésére, a kiberbiztonságban pedig behatolás detektálásra használják.
- **Felügyelet nélküli tanulás (Unsupervised Learning):** Rejtett mintázatokat keres címkeztetlen adatokban. Tipikus alkalmazása a marketingben a fogyasztói szegmentáció (klaszterezés), illetve a hálózati anomáliák felismerése.
- **Megerősítéssel tanulás (Reinforcement Learning):** Egy ágens próbálkozások és jutalmak útján tanul. Ez a módszer kritikus az **önvezető járművek** trajektóriatervezésében és a komplex játéktaktikák kidolgozásában. [1] [4]

Ágazatspecifikus alkalmazások összehasonlítása

Tekintettel hogy több tanulmányt vizsgáltunk meg, az alábbi témaösszehasonlítás jött ki eredménynek. A források technikai szempontból hasonlóságot mutatnak de a célok és az alkalmazott modellek eltérők. Közlekedés és autonóm rendszereknél a fókusz a biztonságon és a valós idejű döntéshozatalon van. A mélytanulást (CNN) képfeldolgozásra, a megerősítéssel tanulás, pedig a jármű irányítására használják. Itt a hiba emberi életetekbe kerülhet, szemben a marketinggel.

- **Kiberbiztonság:** A fő feladat az **anomália-detektálás** és a behatolásmegelőzés. A források hangsúlyozzák, hogy a pontosság (accuracy) helyett a **felidézés (recall)** a döntő mutató, mivel egy támadás elvétsége súlyosabb következményekkel jár, mint egy téves riasztás. Itt megjelenik a kvantumszámítógépek jövőbeli szerepe is az észlelés gyorsításában.
- **Marketing és gazdaság:** A tradicionális ML-megoldások (pl. döntési fák) népszerűbbek az **interpretálhatóság** miatt. A cél a fogyasztói élmény fokozása és a profittermelés az adatok bányászata révén.
- **Szociológia és médiaelmélet:** Itt a gépi tanulás egyfajta „**technológiai protézis**”, amely tágítja szellemi kapacitásunkat. A szociológiai alkalmazások legnagyobb kihívása a **humán annotálás (címkézés) szubjektivitása**, mivel a komplex fogalmak (pl. depresszió keretezése) interpretációja még az embereknek is nehéz.

Elméleti és filozófia háttér

A gépi tanulás egyedülálló filozófiai párhuzamokat von, amelyet a források több irányból is kifejtenek. David M. Berry a gépi tanulás elméleti hátterét Baruch Spinoza filozófiájára alapozza. Spinoza két fogalmat alkalmaz a számítási folyamatok leírására:

- **Számító művelet (Natura naturans):** Az aktív tanulási komponens, amelyben a rendszer algoritmusa folyamatosan fejlődik és adaptálódik, hasonlóan a teremtő természethez, amely a világ változásait és fejlődését irányítja.

- **Kiszámított művelet (Natura naturata):** Az elvégzett számítások eredményeként létrejött, passzív állapotban lévő modell, amely végrehajtja a szükséges feladatokat, analógiát képezve a már kialakult természet rendjével. [1] [2][3][4][5][6][7][8]

A fentiekkel szemben Gyarmati Péter más megvilágításba helyezi az MI -t az úgynevezett antropomorf modellt helyezi az előtérbe, amely alapján kimondható, hogy az MI az emberi gondolkodás és döntéshozatal egy utánzása. De különbséget tesz az emberi agy, és a gépek közötti működésnek.

Gyarmati a következőképpen gondolkodik a meglátásom alapján, az emberi agy az összefüggéseket felismeri, új célokat tűz ki, amelyeket cselekvéssel megvalósíthat. Ha viszont a számítógép alapján vizsgálódunk, tudhatjuk, hogy a gép egy parancskövető, amit beprogramoznak egy szellemi alkotáson keresztül, egy szoftverben megvalósított funkcióba ágyazzák és kódolt instrukciók alapján feladatot hajtanak végre. Ezzel szemben az emberi agy a születéskor a kinti világból fogadja be az információkat, és idővel megtanulja az anyanyelvét is.

A források rávilágítanak, hogy míg technikai alapok hasonlóak, de az alkalmazott modellek eltérőek lehetnek. Ha közlekedés biztonságot vizsgáljuk itt ebben az esetben a forráscikkek szerint, a fókusz a biztonság és a valós idejű döntéshozatalban van, a mély tanulást a CNN képfeldolgozásra, a megerősítéses tanulást pedig az adott jármű irányítására használják. A kiberbiztonság esetében van közös pont a közlekedés biztonsággal mert itt is van döntéshozatal, de a fő feladat az anomália detektálás, és behatolásmegelőzésben van. A források hangsúlyozzák hogy a felidézés a döntő mutató. [7]

A források mélyebb elméleti kérdéseket is adnak a gépi tanuláshoz, David M. Berry Baruch Spinoza terminológiáját hívja segítségül: a **számító műveletet** (aktív tanulás) a *Natura naturans*-hoz, a **kiszámított műveletet** (passzív, generált adat) pedig a *Natura naturata*-hoz hasonlítja, Gyarmati Péter ezzel szemben az **antropomorf modellt** emeli ki, rámutatva, hogy az MI az emberi gondolkodást igyekszik utánozni, de a számítógép imperatív (parancskövető) működése alapvetően eltér, az emberi agy deklaratív természetétől. Tehát az emberi agy egy deklaratív, mert összefüggéseket tár fel a célok mellett, míg a számítógép imperatív mert parancskövető, és egy beprogramozott utasításkészletet hajt végre. Az alábbi táblázat jól összefoglalja a felügyelt, felügyelet nélküli és a megerősítéses tanulás jellemzőit. [7][9]

Szempont	Felügyelt tanulás	Felügyelet nélküli tanulás	Megerősítéses tanulás
Adatok jellege	Címkézett adatok (példahalmazok)	Címkézetlen adatok	Környezeti visszacsatolás (jutalom/büntetés)
Cél	Klasszifikáció (osztályozás) vagy regresszió	Mintázatok, klaszterek keresése	Optimális stratégia megtalálása próbálkozásokkal
Alkalmazási példa	Lemorzsolódás előrejelzés marketingben	Vásárlói szegmentáció, anomália detekció	Önvezető autó trajektória tervezése

1 táblázat: Tanulási jellemzők, saját szerkesztés.

KÖZÖS KIHÍVÁSOK, TORZÍTÁS ÉS ETIKA

A gépi tanulás alkalmazásával kapcsolatos kihívások fontos szerepet játszanak a mindennapi életben [1][2]. Az algoritmusok és azok adathalmazok torzítást eredményezhetnek mivel a gépi tanulás módszerei gyakran adhatnak más eredmény kimeneteket. [3][4]

Az MI torzításra tekintve az algoritmusok képesek lehetnek arra, hogy valamilyen szempont szerint kiválogassák egy adott állásra a kívánt személyeket, akik megfelelnek a feltételeknek, és azon belül is válogathatnak. Pl. Egy toborzó algoritmus program, amely akár az életkor alapján is válogathat, ami diszkriminációs jelenséget mutathat. Ez a probléma arra hívja fel a figyelmet, hogy az algoritmust finomítani kell, és optimálisan működést biztosítani az adathalmazon, de a fenti problémára tekintve a társadalmi minták érvényesülhetnek [3]. A másik példa az önvezető autók kérdése, az etikai dilemmák egy esetleges bekövetkezett baleset után komoly vitákat válthat ki pl. ki a felelős a balesetért, ki volt a hibás. A médiaelméleti megközelítés pedig arra mutat rá, hogy az algoritmusok nem csupán technikai entitások, hanem egyben ideológiai struktúrák is, amelyek a társadalom értékeit és előítéleteit tükrözik [6]. A gépi tanulás nemcsak egy technológiai eszköz, hanem az emberi tudás kiterjesztése, amelyet alkalmazhatnak a különböző tudományágban.

KIBERBIZTONSÁG

Közlekedés és autonóm rendszerek

A közlekedési rendszerekben a mesterséges intelligenciának a fő célja a biztonságos közlekedés növelése, kiemelve, az önvezető járművek úgynevezett autonóm rendszerek tervezése és fejlesztése során. A gépi tanulás kulcsszerepet játszik a járművek trajektóriájának optimalizálásában. Az autó a teszt pályán a tanulás során képes lehet figyelembe venni a forgalmi adatokat, a környezeti tényezőket, a jármű irányításának hatékonysága céljából, de a közlekedési jelzések, és jelek felismerése kihívást jelenthetnek.

A kiberbiztonság pl. egy informatikai rendszer behatolásában szintén a gépi tanulás kerül az előtérbe az anomáliák felismerése szempontjából. A szignatúra alapú detektálás, jól működhet az ismert támadások ellen. Azonban a forrásművek alapján elmondható, hogy a gépi tanulás a mélytanulás alkalmazása erősen ágazatfüggő, és az élet minden területén eltérőek lehetnek a célok, módszerek és a megoldandó problémák. A közlekedésben és az autonóm rendszerekben az MI elsődleges feladata a biztonság növelése, ahol a képfeldolgozás, mint képi információk elemzése, valamint a megerősítéses tanulás és az MI döntése áll a fókuszban. A kiberbiztonság területén a gépi tanulás az anomáliák felismerésében nélkülözhetetlen szerepe van, az észlelési folyamatok mellett. A közeljövőben a kvantumszámítógépek segítségével gyorsítani lehet az észlelési folyamatokon, de nem utolsósorban a párhuzamos programozás módszerének alkalmazása is segíthet a folyamatok gyorsabb feldolgozására. Ha a marketingkutatásra gondolunk az MI -t a gazdasági elemzésekben kérhetjük és alkalmazhatjuk, amely által jobban megérthetjük az üzleti folyamatokat, és annak a döntéseit. Itt előnyös, ha alkalmazzuk a döntési fákat, amelyek alkalmasak az előrejelzésekre is. A mesterséges neuronhálók és a mély tanulás egy adott struktúrára épül fel, és az úgynevezett backpropagation algoritmussal tanul. A több réteggel rendelkező mély hálózatok hatékonyak a beszéd, nyelv, és a képfeldolgozásban az RNN és LSTM modellek használata során. Összességében a gépi tanulás egy sokoldalú eszköztár, amelynek az alkalmazhatósága, az adatok és a alkalmazása területen való illeszkedést jelenti. Én úgy gondolom,

hogy a jövő kulcsa az MI modellek készítésében, és a hozzá való adathalmaz alkalmazásban van. Az alábbi mélyebb összehasonlítás öt fő dimenzió mentén mutatja be a források közötti összefüggéseket és eltéréseket.

Az elemzett források alapján a gépi tanulás nemcsak egy technikai módszertan, hanem filozófiai módszertan szempontból is vizsgálándó. Eltérően értelmezik az MI természetét míg David M. Berry Spinozára építve a gépi tanulást generatív, önmagát újratermelő folyamatként írja le, addig Gyarmati Péter hangsúlyozza, hogy az MI alapvetően imperatív, parancskövető rendszer, amely nem képes önálló célalkotásra, bár én ezt vitatom hiszem az MI öntanulásra is képes lehet. A médiaelméleti irányzat e kettő között helyezkedik el, az MI-t az emberi gondolkodást kiterjesztő technológiai protézisként értelmezve. Ha a módszerek alapján vizsgálódunk a forrásanyagok elismerik a felügyelt és felügyelet nélküli, valamint a megerősítéses tanulás szerepét, de az egyes ágazatok eltérő modellek és algoritmusokkal és eltérő adathalmazokkal dolgozhat. A közlekedésben a megerősítéses tanulás dominál, az autonóm döntéshozatala miatt, azonban a kiberbiztonságban pl. számítógépes hálózatban a hibrid és anomáliaalapú megoldások lehetnek a hatékonyak, A mérnöki tudományágban a fuzzy logika ad kapcsolódási pontot, az emberi érzékelés és a matematika között. Az adatok minősége minden területen kiemelt fontosságú. A műszaki és informatikai tudományágban többnyire mérhető adatokra támaszkodik. Az MI esetében a torzított adatok torz döntések eredménye lesz, a felelőség az emberé. A szerzők az alábbiak alapján gondolkodtak az elemzett forrásmunkájuk alapján. [2] [3][4]

- Az MI természete gondolkodik-e a gép?

A gépi tanulás kvázi „teremtő” folyamat, amely belső mintázatokat épít fel, ezért nem pusztán mechanikus. Az MI nem gondolkodik, csak végrehajt. Az antropomorf értelmezés félrevezető, mert elfedi az emberi felelősséget.

- **Ellentét lényege:** generatív autonómia vs. eszközjelleg.

A „fekete doboz” elfogadható ár a magas prediktív teljesítményért (pl. autonóm közlekedés)

- **Társadalomtudományi források:**
Az értelmezhetőség fontosabb, mint a maximális pontosság.
Ellentét: hatékonyság ↔ magyarázhatóság.

Hatékonyság szerinti összevetés ((kontextusfüggően)

- **Leghatékonyabb megközelítések:**
anomália-detektálás + autoencoder
→ ritka, de veszélyes események felismerésében kiemelkedő.
megerősítéses tanulás
→ dinamikus, valós idejű döntésekhez ideális.
- **Közepesen hatékony, de stabil megoldások:**
 - **Marketing:** döntési fák, logisztikus regresszió
→ jó kompromisszum predikció és értelmezhetőség között.
 - **Mérnöki alkalmazások:** fuzzy logika
→ robusztus, de korlátozott adaptivitás

- A források közötti viták arra mutatnak rá hogy a gépi tanulás nem univerzálisan jó vagy rossz. A hatékonyság mindig attól függ hogy
- Milyen a vizsgált probléma mire keressük a megoldást
- Milyen MI modellt alkalmazunk vagy milyen saját modellt fejlesztünk
- Milyen adathalmazt adunk meg a modellnek
- És milyen következtetést vonunk le a modell tanításából.

FŐBB ELLENÉRVEK ÉS SZEMLÉLETBELI ÜTKÖZÉSEK

Míg a mérnöki és informatikai források alapján (Horváth Gábor, Husi Géza), a neurális hálózatokat az emberi és biológiai modellként kezelik, amely az idegsejt működését utánozza, addig Gyarmati éles kritikát fogalmaz meg ezzel szemben, ő azt mondja hogy a számítógép működése imperatív azaz parancskövető míg az emberi gondolkodás deklaratív ami azt jelenti hogy az összefüggéseket tárja fel, Gyarmati úgy véli hogy a gép üres agyal nem képes önálló célok megvalósítására, csak szoftveres tudás és ismeret feltöltéssel, így az MI soha nem lesz egyenrangú a humán intelligenciával. A marketingforrásra tekintve Pál és Iványi hangsúlyozzák hogy az üzleti döntéseknél a tradicionális döntési fák hatékonyabb lehet, mert az eredmények jól magyarázhatóak, de ezzel szemben a kiberbiztonság esetében a mély tanulás (Deep Learning) részesítik előnyben mert bonyolult mintázatokat is képes felismerni. A műszaki cikkek az adatokat egzakt mérési eredményeknek tekintik.

Németh András szerint a felügyelt tanulás alapját adó címkézésben nincs igazság, azaz ez szubjektív és interszubjektív, az algoritmus pedig átveszi az annotátorok előítéleteit. Az annotálás itt azt jelenti hogy a tanulóhalmaz létrehozása a kódolók munkáját igényli, ami egy szubjektív és interszubjektív folyamat.[6]

HATÉKONYSÁG

A források összehasonlítása alapján megállapítható, hogy a hatékonyság attól függ, hogy milyen feladatra, és milyen területen alkalmazzuk az MI-t. Minden területen más és más algoritmus bizonyulhat hatékonyaknak.[1] Nem létezik univerzális módszer, de egyes problémátípusra beazonosítható a gyakorlat alapján a megfelelő algoritmus alkalmazása, a gépi modellben. Ha egy önjáró autó programjának a fejlesztését végezzük akkor az útkeresés és optimalizálás feladatokban az A* kereső algoritmus adhat hatékonyabb működést a mohó algoritmussal szemben, mert nem csak a becsült távolságot vizsgálja hanem az útvonal költségeket is, ami az optimális megoldást adhatja, és meghatározza a legrövidebb utat is hatékonyabban, elhárítja hogy az algoritmus végtelen utakon fusson.[5][6] A játékprogramok esetében a DeepMind AlphaGo Zero amely a megerősítéses tanulás hatékonyságát adja, mert ez az algoritmus bizonyítja hogy komplex döntésekben is képes az optimális stratégiák kiváltására. [10][11] A beszédfelismerésben az LSTM hálózatok bizonyultak a leghatékonyabbnak, mert képes az összefüggések megőrzésére, amely fontos szempont a beszéd értelmezésében, amely a szavakban és mondatokban rejlik a jelentés. Összességében a források azt támasztják alá, hogy a hatékonyság nem kategóriának tekinthető, hanem az adott problémát vizsgálva alkalmazzuk a megfelelő algoritmust, amelyik jól teljesít a megoldásban. A szignatúra alapú detektálás ez a legkevésbé hatékony a támadások ellen, mivel csak a már ismert mintákat ismeri fel. A torzított tanulóhalmaz esetében a források egyet-

értenek, hogy bármilyen ML modell hatástalan lehet, ha az adathalmaz torzítottak. A források több konkrét modellt is kiemelnek, amely kiemelkedő hatékonyságot mutatnak. Ilyen például a V3 hibrid modell, amely Brunner Csaba kutatása alapján. Ez a modell a leghatékonyabb a ritka és komplex támadások detektálásában, míg a többi modell elvétheti az alul prezentált osztályokat, a V3 mas modell kb 65,68% os felidézési arányt tudhat magának (recall arány) ért el, ami messze meghaladja a hatékonyság százalékát a hagyományos más modellekkel szemben. De van egy másik modell, a (CVAE-DNN) a szakirodalmi összehasonlítások alapján ez a modell (javított kondicionális variációs autoencoder és mély neurális hálózat kombinációja) kimagasló, **85,97%-os pontosságot (accuracy)** és **62,66%-os felidézést** mutatott az NSL-KDD adathalmazon.

Az RNN (Rekurrens Neurális hálózat a mélytanulási modellek közül az egyik leghatékonyabbnak mondható a források alapján, mert a hálózati forgalmat időbeli mintázatainak felismerésére képes. A leghatékonyabb modellek (mint a fent említett V3) képesek a többségi osztályokon elért teljesítményt feláldozni azért, hogy a veszélyes támadásokat nagyobb arányban ismerjék fel, azonban van lehetőség a szintetikus mintavételre (SMOTE) és mesterséges minták generálásával segíteni az algoritmust (SVM SOTE) a támadások felismerésében. A detektálási pontosságot pedig a TPE (Tree-structured Parzen estimators) algoritmussal jelentősen lehet javítani a detektálási pontosságot. A hálózati anomáliák észlelését, a kvantumszámítógépek megjelenésével gyorsíthatóak. [3]

Az észlelési képesség javítható a hálózatban, ha a hibrid modellt alkalmazzuk, ami szignatúra alapú (ismert mintákat kereső), és az anomális alapú (nem szokványos viselkedés a hálózati forgalomban) detektálásnak kombinációjának alkalmazzuk. Ez különösen jól alkalmazható a mély autoencoder hálózatoknál, ahol a normál forgalom tanulása mellett képesek felismerni az eltérő anomáliákat is. A modellek mellett többféle mérőszámot alkalmazhatnak, amely az adott modell teljesítményét adja vissza, amelyet matematikai és statisztikai mutatószámokkal támasztanak alá. Az alábbi táblázatban összefoglaltam a teljesítmény mutatókat.

A felidezés, detektálási ráta az egyik legfontosabb mutató a hálózati behatolás esetében, ez a támadások arányát méri az összes támadáshoz képest. Ezt azért veszik prioritásként mert egy valódi támadás súlyosabb következménnyel jár, mint egy téves riasztás. A jól osztályozott esetek aránya az összes esethez képest a pontosságot mutatja. A konfúziós mátrix pedig abban adhat segítséget, hogy ezt szemlélteti, a (TO) valódi pozitív, (TB) valódi negatív, téves pozitív (FP), téves negatív (FN) döntéseknek a számát, így a modellt részletesebb elemzés alá tudjuk vonni. A hálózati támadást kategóriaként is mérhetik mivel egyes modellek jobban teljesíthetnek a támadás felismerése során. Téves riasztás esetében. [8]

Teljesítmény mutatók

A teljesítményi mutatók kiemelten fontosak egy kiberbiztonsági rendszer esetében. A recall itt azt méri, hogy milyen arányban méri fel a támadásokat. Az adott rendszer biztonsága esetén ez az egyik kiemelt mutatószám, mivel egy támadás védelem igen komoly következménnyel járhat, ami a teljes rendszer leállítását is okozhatja az adatvesztés mellett. Az Accuracy (pontosság) azaz itt az összes esetet vizsgáljuk és ahhoz mérten lesz a helyes aránymutató. Azonban vigyázzunk mert normál hálózati forgalom nagyobb arányban lehet egy támadást tekintve és más eredményt kaphatunk a környezetben. A konfúziós mátrix a modellben az algoritmus lépéseit bontjuk le döntéseit bontja le kimenetekre, ami lehet

valódi pozitív (TP), valódi negatív (TN), hamis pozitív (FP) és hamis negatív (FN) esetek. Ezek segíthetnek a biztonsági kockázat elemzésében, érdemes az kimenetek értelmezhetőségét segíteni recall vagy precision alkalmazásával kombinálni. [2] [8]

Tesztelési módszertan

Mutató	Mit mér?	Miért fontos a biztonságban	Kockázatok
Recall (Felidézés)	A helyesen felismert támadások aránya az összes tényleges támadáshoz képest	Kritikus, mert egy elmulasztott támadás (FN) súlyos következményekkel járhat	Magas recall gyakran több téves riasztással jár
Accuracy (Pontosság)	Helyes besorolások aránya az összes esethez képest	Általános összehasonlításra használható	Kiegyensúlyozatlan adatoknál félrevezető
Konfúziós mátrix	TP, TN, FP, FN értékek megoszlása	Részletes képet ad a modell hibáiról	Önállóan nehéz értelmezni aggregált mutatók nélkül

2. táblázat: Teljesítmény mutatók (saját szerkesztés).

Konfúziós mátrix

A modellek hatékonysága a konfúziós mátrix segítségével jellemezhető. Mérti úgy tudunk a konfúziós mátrixal hogy lebontjuk a fentiekben leírtak alapján, azaz valódi pozitív (TP), valódi negatív (TN), Hamis pozitív (FP), Valódi negatív (FN). A valódi pozitív (TP) a modell helyesen felismerte a vizsgált eseményt. A valódi negatív esetében a modell helyesen beazonosította a negatív esetet. A hamis pozitív azt jelenti, hogy téves riasztás történt, a modell olyan esetben jelzett téves riasztást amikor normál működés volt. A hamis negatív azt jelenti, hogy a modell a valós eseményt nem vette figyelembe, amit tévesen normál működését osztályozta. A konfúziós mátrix eredményekből statisztikai mutatói kiértékelések hozhatók létre, amelyből a modell jósága meghatározható a kívánt alkalmazási területtől függően.

Tesztelési módszertan

Az adatkezelési módszerek befolyásolják a mért teljesítményi értékeket. Itt van két adat az egyik a tesztadat, a másik a tanító adathalmaz. A tanító adatokon dolgozik a modell, a teszt adatokon pedig az értékelések kimenete várható. Ez segíthet azon, hogy a modell csak az adatokon tanuljon és az ismeretlen adatokon. A KDD Cup 1999 vagy az NSL-KDD alkalmazása lehetővé teszi az eredmények hasonlóságának vizsgálatát más kutatási eredményekkel, így a hatékonyságokból levonhatóak könnyebben a következtetések. Az információszivárgást azt kerülnünk el, mert vannak előfeldolgozási lépések, ami a tanító adathalmaz statisztikai adatokra kell épülniük, ellenkező esetben más eredményt kaphatunk. A osztályszintű vizsgálat során a támadástípusokat ajánlatos külön megvizsgálni, mert az összetettségben a mutatókban rejtett maradhat a támadások felismerése ami igen fontos biztonsági kockázatot is jelenthet. Az alábbi 3. táblázat bemutatja a szempontok alapján a tesztelési módszertani jellemzőket az adathalmazon. [2]

Szempont	Leírás	Jelentőség
Tanító–teszt szétválasztás	A modell csak tanító adatokon tanul	Megakadályozza a túlillesztést
Standard adatkészletek	KDD Cup 1999, NSL-KDD	Összehasonlíthatóság a kutatások között
Információszivárgás elkerülése	Normalizálás, skálázás csak tanító adatok alapján	Reális teljesítménymérés
Osztályszintű értékelés	DoS, Probe, R2L, U2R kategóriák külön vizsgálata	Ritka, de kritikus támadások felismerése

3. táblázat: Tesztelési módszertan és adathalmazok, saját szerkesztés.

Kutatásom alapján a következő eredményekre jutottam a hagyományos és a hibrid modellek szempontjából.

Hibrid modellek	
Modell	Pontosság
V1(Stacking Neurális hálózat)	78,11 % pontosság és 55,84% felidézés(recall) ért el
V2 (Stacking NN + SMOTE Tomek)	78,34%-os pontossággal és 59,29%-os felidézéssel zárt
V3 (Autoencoder + Stacking NN)	A pontossága alacsonyabb (74,26%), a kiberbiztonságban kritikus felidézési aránya (recall) kimagasló, 65,82% volt, különösen a ritka támadások (U2R) észlelésében

4. táblázat: Hibrid modellek összehasonlító mérések, saját szerkesztés.

Hagyományos modellek	
Modell	Pontosság
KNN (K-legközelebbi szomszéd)	76,51% pontosság, 48,3% felidézés
RF (Random Forest)	76,49% pontosság, 48,84% felidézés
DNN (Mély Neurális Hálózat)	80,22% pontosság, 52,77% felidézés
SVM (Support Vector Machine)	72,28% pontosság, 45,88% felidézés

5 táblázat: Hagományos modellek összehasonlító mérések, saját szerkesztés.

Az eredmények összevetésekkor a források hangsúlyozzák, hogy a recall(felidézés) mutató a legfontosabb, mert a ritka és veszélyes támadásokat felismerése kritikusabb a rendszere biztonsága szempontjából. A standarnizált adathalmazok használata szükséges mert az ad valós képet és eredményt a modellek teljesítményéről.

Gyakorlati szempontok és korlátok

A rendszerbe való behatolás esetén több tényezőt kell figyelembe venni, tekintettel az alábbiakban összefoglalt tényezők táblázatára, amelyben láthatjuk, az előnyök és hátrányokat. Az anomália alapú detektálás képes felismerni a null napi támadásokat, ami ismeretlen és a rendszer viselkedését figyelemmel kíséri ez a detektálás, azonban a nem szokványos viselkedést támadásnak veheti és így azonosíthatja be. A szignatúra alapú módszerek esetében az előre definiált támadási mintára épül, de az ismeretlen támadásokkal szemben

gyenge teljesítményű eredményt hozhatnak, azonban, ha ismert a támadási jellemzők akkor stabil működést eredményezhet. A tanítási és futási idő igen fontos tényező lehet egy rendszerben. Azok a modellek, amelyek nagyobb komplexitásúak számottevő a számítási és erőforrásigényük, ami korlátozottságot jelenthet, de a tanítási időt is tekintve a gyors reagálás fontos. A humán autentikáció során lehetőségünk van arra, hogy a modellbe beépítsük a szakértői tudást, ami által lehetőségünk van a modell pontosságának a működésére.

Tényező	Előny	Hátrány
Anomália-alapú detektálás	Új (null-napi) t támadások felismerése	Magas téves riasztási ráta
Szignatúra-alapú módszerek	alacsony FP arány	Ismeretlen támadásokkal szemben gyenge
Tanítási / futási idő	Gyors reagálás kritikus rendszereknél	Komplex modellek erőforrás-igényesek
Humán annotáció	Szakértői tudás beépítése	Szubjektivitás, címkézési hibák

6. táblázat: Gyakorlati szempontok és korlátok, saját szerkesztés.

Algoritmusok hatékonysága

A kiberbiztonságban számtalan algoritmus létezhet, amelyet az alábbi táblázat példaképpen mutat. Az egyes alkalmazási területek eltérő probléma szerkezetek és hatékonysággal rendelkezhetnek, így nem igazán létezik univerzális algoritmus, meg kell nézni milyen területre szeretnénk az alkalmazást. A megfelelő módszer kiválasztása mindig az adott feladat céljától, a rendelkezésre álló adathalmaztól és az elvárt teljesítménytől függ.

Terület	Preferált algoritmus	Miért ez a leghatékonyabb?
Kiberbiztonság	V3 Hibrid (AE+Stacking)	Jobban felismeri a ritka, kritikus támadásokat (Recall).
Közlekedés	DDPG (Reinforcement)	Emberi tudás nélkül is optimális trajektóriát tervez.
Keresés	(A-csillag)*	Optimális és teljes; figyelembe veszi a megtett költséget
Marketing	Tradicionalis ML	Az üzleti döntésekhez szükséges interpretálhatóságot nyújtja
Mérnöki munka	Fuzzy logika	Képes kezelni a pontatlan, emberi jellegű bemeneteket.

7. táblázat: Algoritmusok hatékonysága, saját szerkesztés.

A kiberbiztonság területén a V3 hibrid modell alkalmazása indokolt lehet, mert hatékonyan ismerheti fel a ritka és kritikus támadásokat, A magas recall érték fontos lehet, mert egy elmulasztott incidens, nagyobb károkat okozhat. A közlekedési rendszerben első körben az önjáró autóra gondolhatunk, itt a DDPG típusú megerősítéssel tanulási algoritmus előnye, hogy az emberi szabályrendszer nélkül is képes optimális döntéseket hozni. A keresési problémákra tekintve az (A-csillag)* látszik a leghatékonyabbnak mert az adott feladatra optimális megoldást ad, és az útvonal költségeket is figyelembe veszi, amely egy gráfos feladat esetén igen jó választásnak mondható keresés esetén. Az üzleti világban az elemzésre tekintve a hagyományos gépi tanulási módszerek kerülnek előtérbe, a jó teljesít

ményt adva, és meg tudjuk érteni miért hozott a modell döntést az adott szituációra az algoritmus alapján. A mérnöki megközelítés szempontjából ott a fuzzy logika segítségével megmondhatjuk a pontatlan jellegű inputot, ahol matematikai értelemben a modellezés nehézségbe ütközhet. Ilyenkor a MATLAB programot hívhatjuk segítségül megalkotva a fuzzy szabályokat is. [4][2]

A feltöltött források összességében egyetértenek abban, hogy a mély tanulás jelenleg a legerősebb eszköz a mintafelismerés során, pl. a képi orvosi diagnosztika, mint a röntgen ahol a CT felvételek automatikus kiértékelése, nagy mennyiségű strukturálatlan adat feldolgozását igényli. A konvolúciós neurális hálózatok (CNN)képesek a formák, és összetett mintázatok felismerésére. Ezek a hálózatok képesek egy CT felvételen pl daganatra utaló jelek detektálására. Ez egy felügyelt tanulási folyamat, hiszen a szakértők által címkézett adatokon tanulják meg az egészséges és kóros állapotokat. A tanulási folyamat elsődleges célja hogy a téves diagnózisok csökkentése. A mély felügyelt tanulási modellek hatékonyan támogatják az anomália és mintafelismerést, riasztást adva minden a normálistól eltérő struktúrát, rendszerviselkedésre.

Brunner Csaba munkái rámutatnak a **hibrid megoldások** jelentőségére, amelyek a mély tanulást más módszerekkel kombinálva növelik a megbízhatóságot és csökkentik a kritikus hibák esélyét. Ezzel párhuzamosan Gyarmati Péter és Németh megközelítései hangsúlyozzák a **filozófiai és etikai kontroll szükségességét**, különösen olyan területeken, ahol az algoritmikus döntések közvetlen hatással vannak emberi életre.

Összességében a CNN alapú mély tanulás a leghatékonyabb eszköz az orvosi képfeldolgozásban. A gépi tanulás jövője úgy gondolom, hogy nem az orvos helyettesítésében, hanem az orvosi döntéshozatal felelős támogatásában rejlik. A fekete doboz jellegű mélytanulás az egyik legmeghatározóbb jelenkori irány, ez az algoritmusok átláthatóságát és értelmezhetőségét célozza meg. [3] [7]

BEHATOLÁST DETEKTÁLÓ RENDSZEREK

A behatolást detektáló rendszerek az IDS két alapvető irányzatot ad a szignatúra alapú és az anomália detektálást, melyek eltérő működésre épülnek a források alapján. A források szerint ezek nem egymással versenyeznek, hanem eltérő biztonsági igényeket szolgálnak ki. A szignatúra alapú támadás az ismert támadások felismerésére szolgál. A működése szempontjából, hogy a hálózati forgalmat összeveti, egy előre definiált támadási minta adatbázissal és azonosság esetén riaszt. Ennek előnye a magas megbízhatóság, és az alacsony téves riasztás az ismert fenyegetések szempontjából, azonban hátrány a null napi támadásokkal szemben, amihez nem áll a rendelkezésre szignatúra. Ezzel szemben az anomália alapú detektálás a rendszer normál működését modellezi és nem mintákra épít, így minden olyan viselkedést, ami eltér a normális működéstől, a modell alapján gyanúsnak minősít. Nagyobb rugalmasságot ad, mint a szignatúra alapú rendszer, mert itt képesség van az új típusú támadások felismerésére, viszont magasabb téves riasztás lehet mert a rendszer nehezen különbözteti meg, a valódi támadásokat a forgalmi mintáktól. Azonban itt az anomáliák több típusát megkülönbözteti, a környezetfüggő és egy egymással összefüggő anomáliákat. Összességében itt leírható, hogy az IDS szignatúra alapú megbízható védelmet nyújtanak, az ismert fenyegetéssel szemben, addig az anomália alapú esetében a támadások ellen erősebbnek bizonyulnak. A források alapján a legjobb, ha a két módszert kombináljuk, amely egyesíti az alacsonyhiba arányt és a szélesebb felismerés képességet.

A forrás alapján Brunner Csaba is a kettő kombinációját a hibridet javasolja a hatékony működésre tekintve. Disszertációjában négy alapvető hibrid architektúrát azonosít: a párhuzamos detektálást, a szignatúra anomália szekvenciális detektálást, az anomália szignatúra szekvenciális detektálást, valamint az összetett kevert megoldásokat.

Kutatásai alapján a V3 as modellt találta a leghatékonyabbnak. Ez a modell a stacking neurális hálózatot kombinál mély autoencoder hálózatokkal, amit normál hálózati forgalom alapján tanítanak be, amelynek függvényében a fent leírtak alapján is alkalmas az eltérő minták és anomáliák felismerésére. A teljesítményt tovább növeli az **SVM SMOTE szintetikus mintavételezés** a ritka támadások kiegyensúlyozására.

A V3-as modell különösen hatékonynak bizonyult a ritka, de súlyos fenyegetések például az U2R támadások – felismerésében, ahol **65,82%-os átlagos recall értéket** ért el. Brunner szerint ez azért is előnyös mert a kiberbiztonságban egy elmulasztott támadás nagyobb kockázatot rejthet magában, mint egy téves riasztás. [3][7]

Szerzők szerinti vélemény összehasonlítás

Nagy Attila és Husi Géza munkásságában a gépi tanulás (ML) alkalmazása több alapvető ponton is találkozik annak ellenére, hogy a kutatási területük eltérő. A neurális hálóban mindkét szerző a neurális hálót tekinti a gépi tanulás egyik legfontosabb eszközének. Tárgyalják a publikációjukban, hogy a neuronhálózat rétegekből épül fel, és az emberi agy alapján a neuronok működését szimulálja. Mindkét szerző tárgyalja a felügyelet és felügyelet nélküli tanulási módszert. Nagy Attila a hálózati anomáliák észlelésénél a felügyelet nélküli tanulást helyezi előtérbe, míg Dr. Husi Géza az ellenőrzött tanulást a hiba-visszaterjesztési (backpropagation) modell kapcsán mutatja be. Mindketten egyetértettek abban hogy a gépi intelligencia áttörése a 21. században a megjelenő BIG DATA technológia, és a nagy teljesítményű számítógépnek köszönhető. Nagy Attila megjegyzi, hogy az algoritmusok már korábban a rendelkezésre álltak, csak az erőforrások nem voltak elérhetőek. Mindkét szerző módszernek tekinti a gépi tanulást, azaz a rejtett minták és összefüggések megtalálását. Minkét szerző kimondta, hogy ha egy modell jól van betanítva akkor, általánosan elvárható, hogy a tanító halmazban nem szereplő, de új adatokra is helyes választ adjon. Az osztályozást mindketten alkalmazzák, Nagy Attila a hálózati forgalmat osztja káros és hasznos kategóriákra, Dr. Husi Géza pedig például a kézzel írott karakterek felismerését vagy ipari folyamatok állapotait osztályozza. Azonban van eltérés is, Nagy Attila a gépi tanulást a jövő kvantum- számítástechnikájával és a mély tanulással kapcsolja össze a kiberbiztonságra tekintve, addig Husi Géza a gépi tanulást a fuzzy logikával kombinált hibrid rendszerként mutatja be a mérnöki és ipari folyamatokat adott helyzetben való irányításra. Pl. Daruk irányítása. [1] [2][3][4][5][6][7][8]

Tanulási tényezők kiválasztása

Husi Géza a tanulási tényezők kiválasztását kísérleti és matematikai módszerekkel támogatja.

- Nagy Attila megközelítése:
- A felügyelet nélküli tanulás a hálózati anomáliák észleléséhez
- Kvantum neurális hálózat súlytalanítása és kvantumkeresés.

Összegzőként Husi Géza heurisztikus és automatizált keresési stratégiákat, módszerek támogatja a tanulási tényező megtalálását, míg Nagy Attila esetében a hangsúly az

automatikus mintafelismerésen van. A tudományos szintmeghatározása jelen pillanatban, a fekete doboz jellegű mélytanulástól a hibrid rendszer felé mozdul el.

A JÖVŐ TECHNOLÓGIÁI

A források szerint a klasszikus behatolás érzékelők elavulnak, mert lassú és több hibás jelzést adhat. A hatékonyság növelésének új iránya a mély tanulásban lehet (Depp Learning) mert ez lehetővé teszi a hatékonyabb null napi támadások hatékony kezelése a többretegű neurális hálózatban. A jelenlegi álláspont szerint a hibrid mély tanuló modellek a leghatékonyabbak, mert képesek a modellek az ismert támadások beazonosítására, de az ismeretlen anomáliák felismerésére is. A kiberbiztonsági modellek hatékonysága a ritka támadások felismerésénél több technológiai és módszertannal javítható lehet. A mély auto-encoder hálózatok különös hatékonyság mutatnak, ők a hálózat normál forgalmán tanulnak tapasztalnak, és ennek köszönhetően képesek lehetnek felismerni a ritka anomáliákat is. Az egymásra épülő modellek segítségével pedig pontosabb osztályzást érhetünk el egy komplexebb eseményél. A gépi tanulás jövőjét tekintve a kvantum neurális hálózatok QNN kiemelt szerepet kaphatnak, tekintettel a kvantum gép gyorsaságára, és a párhuzamos számítási feladatok elvégzésére, amely hatékonyabban teszi lehetővé az adathalmazok feldolgozását a gépi tanulási modell számára, és a komplexebb mintázatok felismerésére. A mély tanulás továbbfejlesztése új architektúrát is foglalhat magában, és neurális hálózatokkal pontosabb eredményhez juthatunk optimalizálva az algoritmust is. A megerősítéses tanulás a világban lévő ismeretek és tapasztalataim alapján elmondható, hogy a robotikában, az önvezető járművekben, és a döntéshozó feladatokban elterjedt. A hibrid rendszerek, amelyek kombinálják a rendszereket, mint a neurális hálózatokat és genetikai algoritmusokat gyorsabb és stabilabb tanulást eredményezhetnek. Az aditív rendszereknek köszönhetően fontos hogy ne felejtse az MI, amit korábbi időkben szerzett tudást jelent, és amit felhasználhat egy következő feladat megoldásához is, hiszen a környezet egy dinamikusan változó területű is lehet, ahol szintén tapasztalhat és tanulhat az MI. Összességében a gépi tanulás az egyik kiemelt terület számomra is, mert ezen keresztül jobban meg lehet érteni az MI algoritmusok elvi működését is. A jövőt illetően pedig a gyorsabb, és pontosabb intelligens rendszerek fejlesztési irányába halad, amely képesek lehetnek komplexebb problémák megoldására, azonban a véleményem szerint az embernek megszabott korlátok közé kell tenni az MI -t hogy ne veszítse el az ember a kontrollt az MI felett.

ÖSSZEFOGLALÓ

Jelen cikk a szerzők munkájának az összehasonlítási téziseket foglalta magában, amelyek közül kiemeltem a gépi tanulás témakörét. A cikk a gépi tanulás alkalmazásait vizsgálja, a gépi tanulás nemcsak egy eszköz, hanem egy olyan jelenség, ami formálja a tudományos és társadalmi struktúrát. A tanulmány kiemelten elemzi a gépi tanulás témakörét. A szerzők bemutatják a neurális hálózatok ismereteit, tanulási modelleket, a felügyelt és nem felügyelt tanulási módszert, ahol az annotáció minősége meghatározó az algoritmusok pontosságában. A forrásanyagok alapján megállapítható, hogy a kiberbiztonsági rendszerek hatékonysága az alkalmazott módszereken múlik, és nem annyira az algoritmuson, hanem a metrikák és a rendszer beállítások alapján. A közlekedés és behatolás észlelései

különböző rendszereket képvisel, és ez által mások lehetnek az elvárások a működésre tekintve. A recall és az osztályszintű értékelésnek kiemelt szerepe van, de a kritikus támadások felismerése fontos tényező. A konfúziós mátrix jól mutatja a modell hibáinak és tulajdonságainak jellegét. Hogy az eredmények más kutatással összehasonlítható legyen, fontos a jó adatkezelés és tanító és teszt adat elkülönítése egymástól, és az adatszivárgás elkerülése. Ezek megadhatják a mért teljesítmény valós képét. A gyakorlat alapján a hibrid megoldások alkalmazása optimálisabb. Több forrásmunkát összehasonlítva a cikk rávilágít különböző megközelítések formájára, a fókusz a gépi tanulás kutatásában és alkalmazásában van.

Szó esik a közlekedési rendszerekről is, amelynek egyik célja a biztonságos közlekedés. A gépi tanulás kulcsszerepet játszik ebben is, figyelembe véve a környezeti tényezőket. Fontosak a célok, hogy mire kívánjuk alkalmazni a gépi tanulást, és az alapján érdemes választani a megoldási módszereket. Említettem a képfeldolgozás fontosságát a gépi tanulás mellett, amelynek szintén fontos szerepe van az önjáró autók hatékony működésében. A kvantumszámítógépek fontosságát is kiemeltem, amely nagyban segíthet a gyorsabb műveletvégzésben, és a minták felismerésében a modellben az adott környezeti területen. A források technikai szempontból hasonlóságot mutatnak, és közös és ellentétes véleményen alapuló fókusz kiemeltem a cikkben, ebből is látszik, hogy az alkalmazott modellek eltérőek lehetnek. A kiberbiztonság esetén vannak közös pontok, mert ott a döntéshozatal van kiemelve, az anomália detektálás és behatolásmegelőzés tükrében. Az MI természetét eltérően értelmezik David M Berry a gépi tanulást egy önmagát újratermelő folyamatként írja le, míg Gyarmati Péter az MI-t egy imperatív parancskövető rendszernek tekinti, amely nem képes önálló alkotásra.

Jelen cikk a szerzők munkájának összehasonlítását, filozófiai gondolkodásának a leírásának a célja, hogy ki mit gondolt és milyen módszert preferál az MI alapon a gépi tanuláson ami igen fontos kulcsszerepet játszik, egy mesterséges intelligencián alapuló program működésében. Amit gondolok, mint szerző a kutatásom során, hogy eltérő vélemények vannak az MI vel kapcsolatban, azonban az alkalmazás a választott területtől függ, amelyben a modell és az MI algoritmus fog majd dolgozni.

FELHASZNÁLT IRODALOM

- [1] D. M. Berry, „Bevezetés a gépi tanulás médiaelméletébe” Médiakutató, vol. XXV, no. 2, pp. 53–61, 2024, doi: 10.55395/MK.2024.2.4.
- [2] T. Bécsi, S. Aradi, and Á. Fehér, „A gépi tanulás szerepe és hatásai a közlekedésben” Közlekedéstudományi Szemle, vol. 70, no. 1, pp. 54–65, 2020, doi: 10.24228/KTSZ.2020.1.1.
- [3] C. Brunner, „Behatolásdetektálás gépi tanulás által” Ph.D. dissertation, Gazdaságinformatika Doktori Iskola, Budapesti Corvinus Egyetem, Budapest, 2020, doi: 10.14267/phd.2020026.
- [4] A. Nagy, „Hálózati anomáliák detektálása gépi tanulással” *Információbiztonság*, 2022.
- [5] A. Németh and K. Virágh, „Mesterséges intelligencia és haderő – A mesterséges intelligencia területei III. rész” Haditechnika, vol. 56, no. 3, pp. 2–7, 2022, doi: 10.23713/HT.56.3.01.
- [6] G. Husi, „Mesterséges intelligencia alkalmazása”. Budapest: TERC Kereskedelmi és Szolgáltató Kft., 2013.

- [7] P. Gyarmati, „Gondolatok a mesterséges intelligencia, a gépi tanulás kapcsán (III. rész)” *Mesterséges Intelligencia*, vol. 5, no. 2, pp. 25–38, 2023, doi: 10.35406/MI.2023.2.25.
- [8] B. Cs. Pál and T. Iványi, „Gépi tanulás lehetőségei a marketingkutatásban” in **A (marketing) világ megkettőződése** – Egyesület a Marketing Oktatásért és Kutatásért XXX. Nemzetközi Konferenciájának Absztrakt- és Tanulmánykötete, K. Szűcs, P. Putzer, and M. Törőcsik, Eds., Pécs: Pécsi Tudományegyetem Közgazdaságtudományi Kar, 2024, pp. 112–118, doi: 10.62561/EMOK-2024-10.