



ÓBUDAI EGYETEM
BÁNKI DONÁT GÉPÉSZ ÉS
BIZTONSÁGTECHNIKAI MÉRNÖKI KAR

Egy hálózat GondolatAI

Oláh István doktorandusz

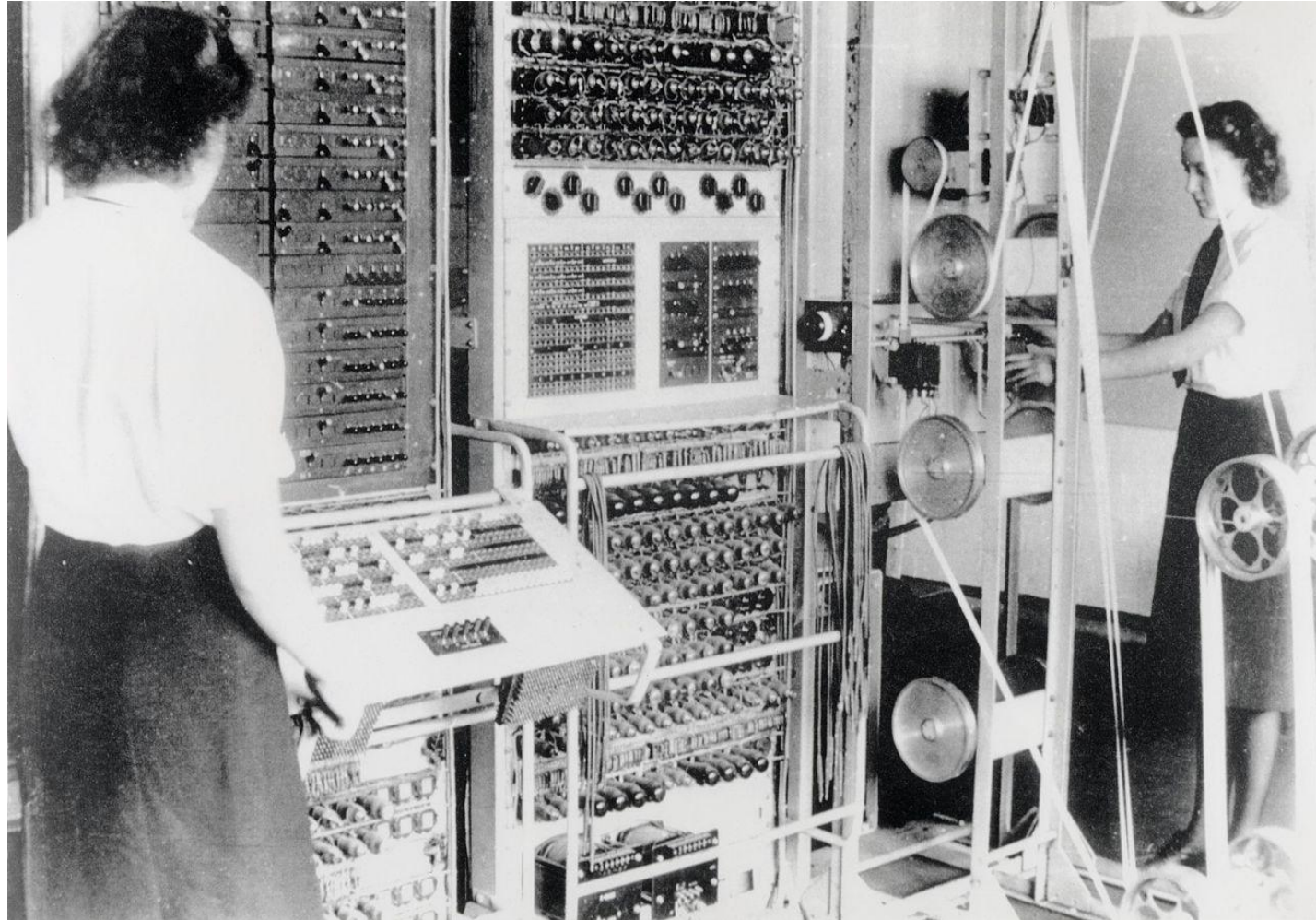
Infokommunikáció 2025 Nemzetközi tudományos-szakmai konferencia

2025. november 13.

NKE Zrínyi Miklós Laktanya és Egyetemi Campus M. Szabó Miklós altábornagy Díszterem, Budapest

A szoftverek (AI) uralják a fizikai eszközöket (embereket?) is !?

- 30 000 éves az első számoló eszköz
- 5.000 éves az abakusz
- 1620 logarléc
- 1623 első számológép
- 1786 TPV=regiszter
- 1812 programozható számlógép, memóriával
- 1820 TPV alapú szövőgép
- 1853 integrált „áramkör” elmélet
- 1887 statisztikai gép USA-népszámlálás
- 1936 elektromechamikus számológép
- 1939 Z1 szabadon programozható gép, No.
- **1943 Colussus Anglia, nem decimális (az AI kezdete)**
- 1944 ENIAC USA
- 1945- Neumann-elvek szerinti gépek USA (előtte Mao?)



- Már az első számítógép is AI machine volt
- 1960 LISP első AI „nyelv” (család)
- **1962 Hazai egyetemeken oktatott tárgyakban benne van az AI**
- 1965 ELIZA-MIT
- 1969 Szakértői rendszerek, MYCIN
- **1985 az első AI appom**
- 1990 ML algoritmusok elérhetőek a civil szférában is, elsőre a fordító appok jöttek
- **1997 Garri Kaszparov-IBM Depp Blue**
- 2006 Deep learning új modellek
- 2011 IBM Watson, nyelvi modellek LLM
- 2014 GAN: Generatív Adeversial Networks
- sok az elérhető adat (de tiszta-e?), olcsó=skálázódó feldolgozási technológiák

- **Artificial Intelligence:** a számítástechnika azon területe, amely olyan intelligens gépek létrehozására törekszik, amelyek képesek az emberi intelligenciát megismételni vagy meghaladni.
- **Machine Learning:** a mesterséges intelligencia olyan részhalmaza, amely lehetővé teszi a gépek számára, hogy a meglévő adatokból tanuljanak, és az adatok alapján döntéseket vagy előrejelzéseket hozzanak.
- **Deep Learning:** olyan gépi tanulási technika, amelyben neurális hálózatok rétegeit használják az adatok feldolgozására és a döntések meghozatalára.
- **Generative AI:** új írott, vizuális és auditív tartalmakat hoz létre kérések vagy meglévő adatok alapján

Az AI szelleme „is” kiszabadul a palackból....



- **Természetes nyelvi chatbotok:** az ügyfélszolgálatától kezdve a személyre szabott korrepetáláson át az öngyilkosság-megelőzési tanácsadók képzéséig minden
- **Fejlett keresési képességek:** nem csak linkek, hanem válaszok keresése
- **Szöveggenerálás:** e-mailek, jogszabálytervezetek, szerződések, forgatókönyvek
- **Tartalom** (hang, videó, szöveg) valós idejű fordítása, átírása, elemzése és összefoglalása
- **Képfelismerés és -generálás.** (Smink felismerés, stb.)

Hack-el, és ... védekezik!

<https://nki.gov.hu/it-biztonsag-kategoria/mesterseges-intelligencia/>

https://nki.gov.hu/wp-content/uploads/2023/03/A-ChatGPT-kiberbiztonsagi-kockazatai_CTI_jelentes.pdf

„A mesterséges intelligencia miatti kipusztulás
veszélyét az olyan emberiségre leselkedő
fenyegetésekkel egyenértékűen kell kezelni, mint
a világjárványok vagy az atomháború.”

NIST AI 100-1

Artificial Intelligence Risk Management
Framework (AI RMF)

National Institute for Standards in Technology (NIST) mesterséges intelligencia-kockázatAI

- 2021-ben a Kongresszus arra utasította a NIST-t, hogy dolgozzon ki egy önkéntes keretet a megbízható mesterséges intelligencia számára
- 2023: Artificial Intelligence Risk Management Framework (AI RMF 1.0)

Célja: olyan iránymutatások és bevált gyakorlatok gyűjteménye, amelyek célja, hogy segítse a szervezeteket a mesterséges intelligencia technológiák bevezetésével és használatával kapcsolatos kockázatok azonosításában, értékelésében és kezelésében

- **Érvényesség és megbízhatóság** – Az AI téves információkat is szolgáltat, amit a GenAI-ban “hallucinációnak” is neveznek. Fontos, hogy a vállalatok validálni tudják, **hogy az általuk alkalmazott mesterséges intelligencia pontos-, és megbízható-e**
- **Biztonságosság** – Annak biztosítása, hogy az információkat ne osszák meg más felhasználókkal, mint a hírhedt Samsung ügyben
- **Ellenállóság és biztonság** – A szervezeteknek biztosítaniuk kell, hogy az AI rendszer védett és biztonságos legyen a támadásokkal szemben és sikeresen meg tudja hiúsítani a kihasználására vagy a támadások segítésére irányuló kísérleteket
- **Átláthatóság** – Fontos, szempont **az AI működésének a megértése**, hisz nincs benne semmiféle fekete mágia
- **Adatvédelem** – Biztosítani kell, hogy a kért információk védettek és anonimizáltak legyenek a felhasználás során
- **Tisztességesség** – A káros előítéletek kezelése (például az AI arcfelismerésben gyakran előfordul elfogultság, a világos bőrű férfiakat pontosabban azonosítják, mint a nőket és a sötétebb bőrszínűeket). Ha például a mesterséges intelligenciát a bűnüldözésben használják, ennek súlyos következményei lehetnek

- A klasszikus ITsec megoldásokat AI képessé kell tenni
- Az AI új kihívásaira új ITsec hálózati megoldások szükségesek

ENISA 2025. októberben publikálta a Threat Landscape 2025.összefoglalóját. Az Európai Unió 2022/2555 számú irányelve (NIS2) 6 szerinti ágazati besorolás szerint a pénzügyi szolgáltatók a hatodik legjobban támadott szolgáltatások a kibertérben jelenleg. A leggyakoribb támadások adathalászattal kezdődnek. Az ENISA szerint ezek

80% már AI alapú.

AI + Large Language Model (LLM)

A
Z

Ú
j

T
O
P

Támadás típusa	Rövid leírás	Főbb példák/hatás	Védekezési stratégia
Prompt Injection	Speciális bemenetekkel manipulált válasz	AI chat szakmai/etikai guardrail megkerülése	Szintaktikus és szemantikus szűrés, input validáció, guardrails, ember-a-hurokban jóváhagyás
Training Data Poisoning	Mérgezett, elfogult vagy rosszindulatú adatokkal tanított modell	Elfogult vagy hibás output, adatvédelmi sérülés	Adatforrás audit, anomaly detection, provenance tracking, adversarial tesztelés
Model Inversion / Membership Inference	Tanítóadatok visszafejtése, PII leak	Személyes adatok kiszivárgása	Differential privacy, likvid lekérdezőszám, titkosítás
Evasion Attack	Finoman módosított inputok, tévesztés céljából	Képfelismerési hibák, azonosítás megkerülése	Adversarial training, input transformation, ensemble modellek
Model Stealing (Extraction)	Modell „lemásolása”, funkcionalitás kisajátítása	IP, tulajdon ellopása	API rate limiting, query log, watermarking
Supply chain attack	Harmadik féltől származó kód, modell vagy dependency támadás	Backdoor/rosszindulatú könyvtár	SBOM, code audit, integrity check, rendszeres függőségfrissítés
Insecure Output Handling	Nem szűrt output, amely downstream rendszerben sérülékenységet okoz	XSS, SQL injection	Output validáció, sandbox, CSP
Hallucination/hamis tartalom	Hihető, de téves vagy veszélyes AI-válasz	Félretájékoztatás, reputációs kár	RAG, fact-check, audit, emberi kontroll
Excessive Agency	AI agent, nem kontrollált autonóm viselkedés	Nem kívánt művelet végrehajtása	Minimális jogosultság, naplózás, felülvizsgálat
Unbounded consumption/denial-of-wallet	Erőforrás (pl. API) túlhasználata/visszaélés	Szolgáltatáskimaradás	Limitálás, throttling, költség-alapú routing

- Hálózati Proxy
- Hálózati IDS
- Hálózati IPS
- Hálózati xDOS védelem
- Hálózati mail védelem
- Hálózati Sandbox-ok

Heurisztika alapú működés

25+ éve

Neurális AI !

- GENAI Proxy
- GENAI IDS
- GENAI IPS
- GENAI xDOS védelem
- GENAI mail védelem
- GENAI Sandbox-ok
- GENAI Tűzfal
- GENAI Router
- GENAI Switch

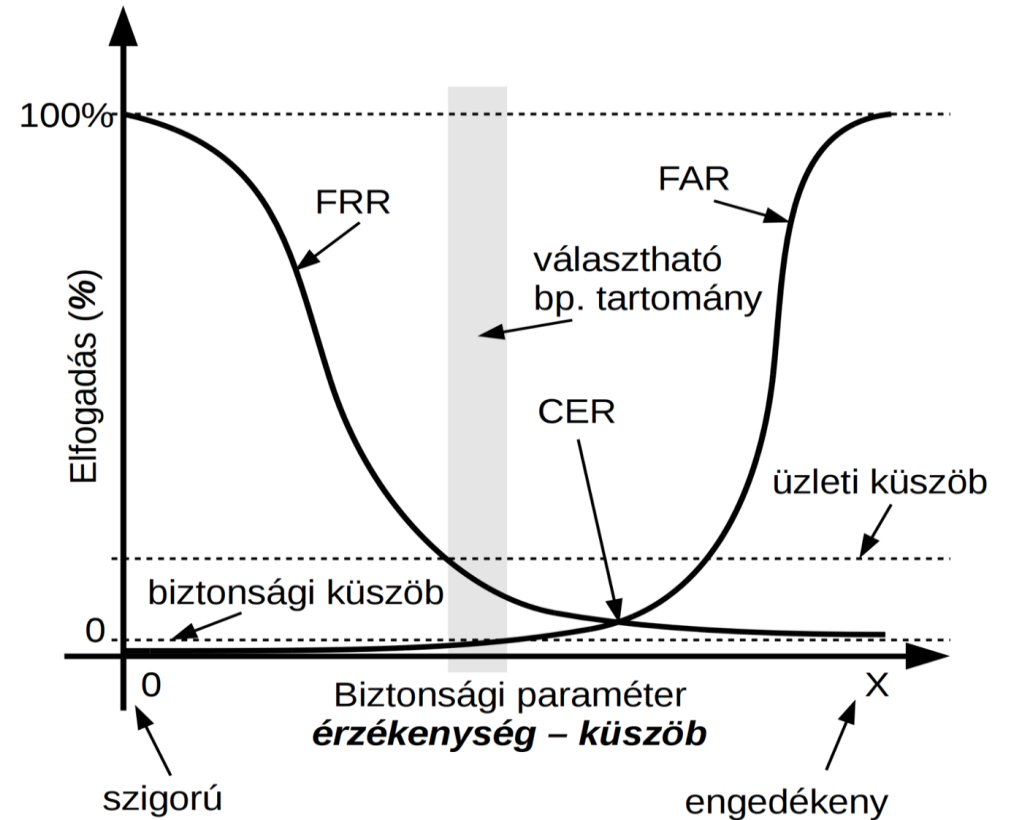
**Már most is
elérhetőek!**

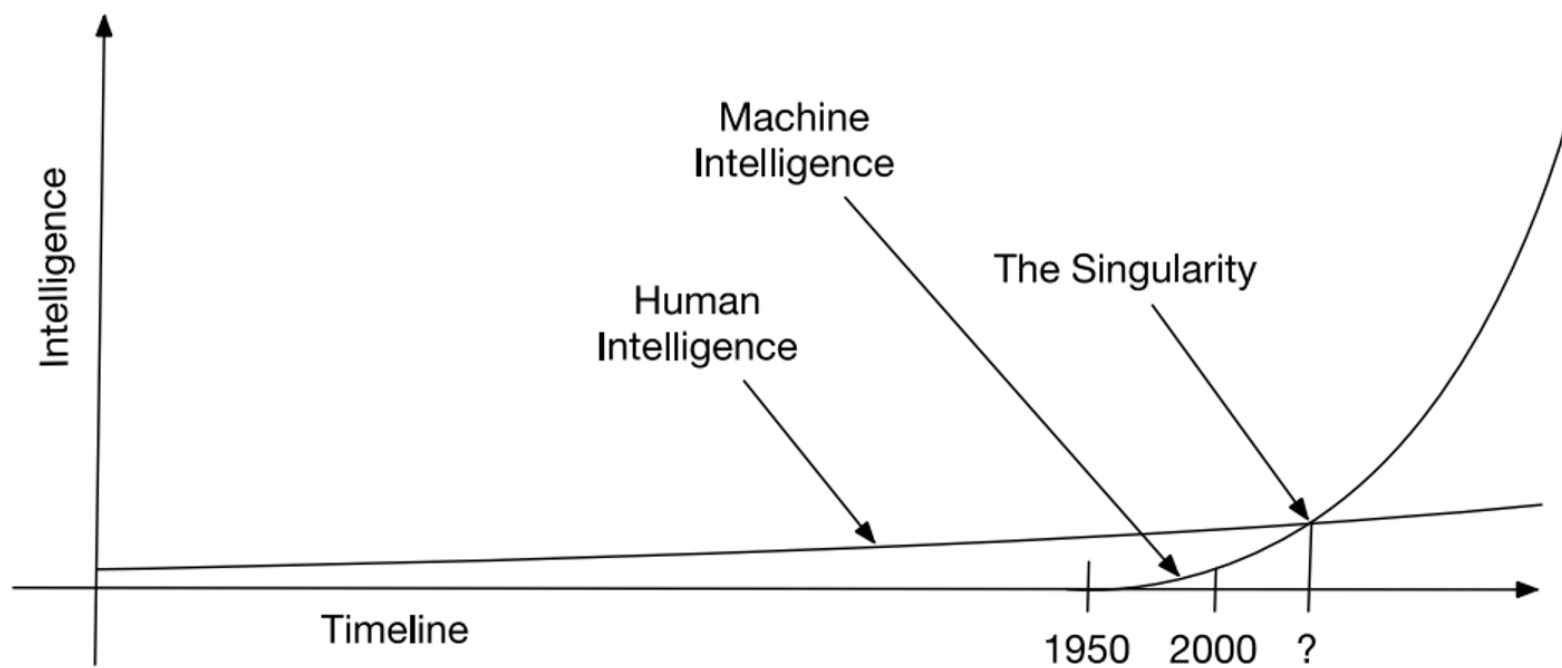
Ez az MI magyarázható működésű, mert MI fejlesztettük,
ezért lehet auditáltatni, és verifikálni is lehet.

A modell úgy tanul, olyan adaton tanul, és azt tanulja amit

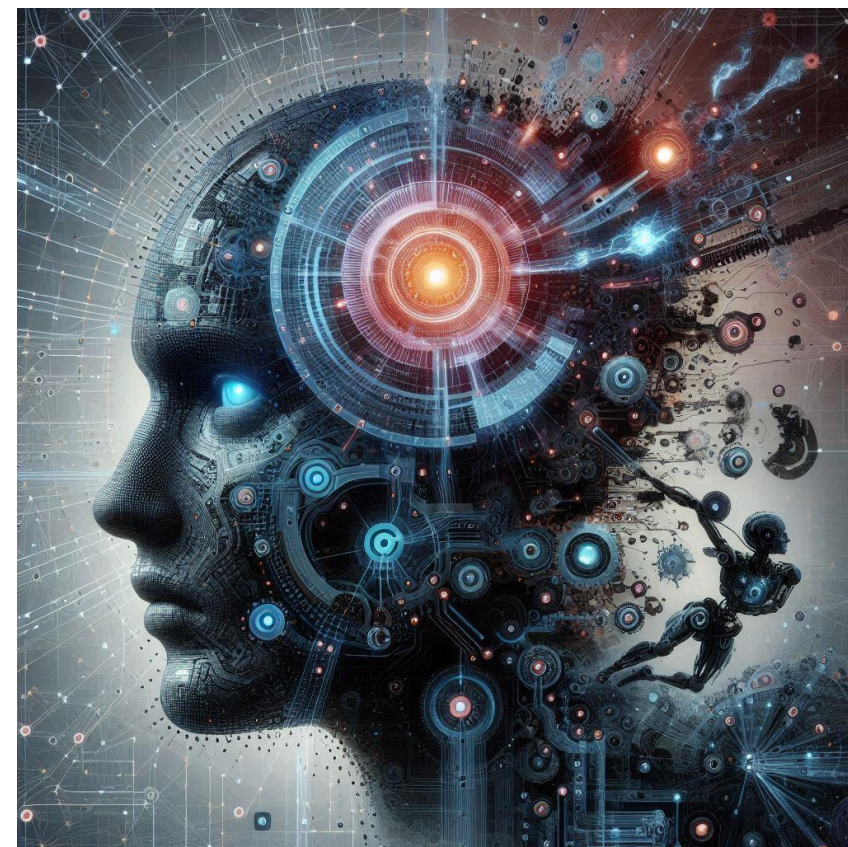
MI szeretnénk, és nem azt amit mások!

- Külső, független szereplő
- Magyarázható, mérhető, megismerhető az eredmény
 - “Érzésre” mégis más – egyedi esetek, egyér befolyása
- A számítások irreverzibilisek
 - Az eredeti aláírás nem állítható vissza!
- Negatív hipotézis + konfidencia intervallum
→ Elvetésre került → 





Forrás: Tóth Márton EIVOK-40



Forrás: Microsoft Copilot

Kvantumbiztonság!

Kvantum+Felhő+AI?

Biológia + Matek processzorok



- 1 AI / LLM Biztonsági Szolgáltatások - KPMG Magyarország. <https://kpmg.com/hu/hu/home/services/advisory/technology/cyber-lab/ai-llm-security.html>
- 2 What Are Adversarial AI Attacks on Machine Learning?. <https://www.paloaltonetworks.com/cyberpedia/what-are-adversarial-attacks-on-AI-Machine-Learning>
- 3 NIST Identifies Types of Cyberattacks That Manipulate Behavior of AI <https://www.nist.gov/news-events/news/2024/01/nist-identifies-types-cyberattacks-manipulate-behavior-ai-systems>
- 4 NIST Trustworthy and Responsible AI NIST AI 100-2e2025 <https://csrc.nist.gov/pubs/ai/100/2/e2025/final>
- 5 ENISA THREAT LANDSCAPE 2025 <https://www.enisa.europa.eu/sites/default/files/2025-10/ENISA%20Threat%20Landscape%202025.pdf>
- 6 Az Európai Parlament és a Tanács (EU) 2022/2555 irányelve (2022. december 14.) az Unió egész területén egységesen magas szintű kiberbiztonságot biztosító intézkedésekről, valamint a 910/2014/EU rendelet és az (EU) 2018/1972 irányelv módosításáról és az (EU) 2016/1148 irányelv hatályon kívül helyezéséről (NIS 2 irányelv) (EGT-vonatkozású szöveg) DIRECTIVE (EU) 2022/2555 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 14 December 2022 on measures for a high common level of cybersecurity across the Union, amending Regulation (EU) No 910/2014 and Directive (EU) 2018/1972, and repealing Directive (EU) 2016/1148 (NIS 2 Directive) <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32022L2555>
- 7 Megjelent a NIST Mesterséges Intelligencia kockázatkezelési keretrendszere <https://nki.gov.hu/it-biztonsag/hirek/megjelent-a-nist-mesterseges-intelligencia-kockazatkzezesi-keretrendszere/>
- 8 ISO/IEC 42001:2023 Information technology — Artificial intelligence — Management system <https://www.iso.org/standard/42001>
- 9 Az Európai Parlament és a Tanács (EU) 2024/1689 rendelete (2024. június 13.) a mesterséges intelligenciára vonatkozó harmonizált szabályok megállapításáról, valamint a 300/2008/EK, a 167/2013/EU, a 168/2013/EU, az (EU) 2018/858, az (EU) 2018/1139 és az (EU) 2019/2144 rendelet, továbbá a 2014/90/EU, az (EU) 2016/797 és az (EU) 2020/1828 irányelv módosításáról (a mesterséges intelligenciáról szóló rendelet) (EGT-vonatkozású szöveg) Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act) (Text with EEA relevance) <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689>
- 10 Bankoknál, biztosítóknál vizsgálta a mesterséges intelligenciát, gépi tanulást az MNB <https://www.mnb.hu/sajtoszoba/sajtokozlemenyek/2023-evi-sajtokozlemenyek/bankoknal-biztositoknal-vizsgalta-a-mesterseges-intelligenciat-gepi-tanulast-az-mnb>
- 11 Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) (Text with EEA relevance) <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32016R0679>
- 12 Survey Shows How Generative AI Puts Organizational AI Maturity to the Test <https://www.gartner.com/en/documents/5373863>
- 13 A (kiber)biztonság gyakorlatAI: Inspiráló előadások és beszélgetések a jövőhöz vezető úton - EIVOK-59. <https://www.hte.hu/informaciobiztonsagi-szakosztaly-eivok/-/esemeny/1/4975549/eivok-59-a-kiber-biztonsag-gyakorlatai>
- 14 Alapú dokumentum hitelesítés a gyakorlatban www.hte.hu/documents/10180/4963359/I-03_Golda+Bence.pdf
- 15 Code security audit based on regulations at a financial institution <https://ieeexplore.ieee.org/document/11030093>
- 16 AI Governance Framework - AI Governance Framework. <https://ai-governance.eu/>
- 17 MULTILAYER FRAMEWORK FOR GOOD CYBERSECURITY PRACTICES FOR AI <https://www.enisa.europa.eu/sites/default/files/publications/Multilayer%20Framework%20for%20Good%20Cybersecurity%20Practices%20for%20AI.pdf>
- 18 NIST AI Risk Management Framework <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>
- 19 REGULATIONS REGULATION (EU) 2022/2554 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 14 December 2022 on digital operational resilience for the financial sector and amending Regulations (EC) No 1060/2009, (EU) No 648/2012, (EU) No 600/2014, (EU) No 909/2014 and (EU) 2016/1011 <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32022R2554>
- 20 13 Az EU AI Act és a GDPR kapcsolódásai <https://dmp.hu/adatvedelem/az-eu-ai-act-es-a-gdpr-kapcsolodasai/>
- 21 Megfelelés az AI Act-nek - gyakorlati példák az adatvédelmi feladatokra. <https://www.jogiforum.hu/blog-adatvedelem-10/2024/10/11/megfeleles-az-ai-act-nek-gyakorlati-peldak-az-adatvedelmi-feladatokra/>
- 22 A generatív AI adatvédelmi kockázatai – Jogászvilág. <https://jogaszvilag.hu/vilagjogasz/a-generativ-ai-adatvedelmi-kockazatai/>
- 23 How to Scale Your Finance AI Pilots <https://www.gartner.com/en/documents/6285783>
- 24 OWASP Top 10 for Large Language Model Applications <https://owasp.org/www-project-top-10-for-large-language-model-applications>
- 25 Navigate threats to AI systems through real-world insights <https://atlas.mitre.org/>
- 26 ML10:2023 Model Poisoning - OWASP Foundation. https://owasp.org/www-project-machine-learning-security-top-10/docs/ML10_2023-Model_Poisoning
- 27 Build a Zero Trust Framework for Secure AI Implementation. <https://www.microsoft.com/en-us/microsoft-365/business-insights-ideas/resources/build-a-secure-zero-trust-secure-foundation-for-ai>
- 28 How is AI Strengthening Zero Trust? | CSA. <https://cloudsecurityalliance.org/blog/2025/02/27/how-is-ai-strengthening-zero-trust>
- 29 Take Back Control and Trust Your Enterprise AI with ZeroTrusted.ai's AI Firewall / AI Gateway / AI Health Check [ZeroTrusted.ai. https://www.zerotrustrusted.ai/](https://www.zerotrustrusted.ai/)
- 30 Hogyan segít kezelni az AI az ellátási láncban jelentkező kockázatokat <https://logisztika.hu/2024/10/14/hogyan-segit-kezelni-az-ai-az-ellatasi-lancban-jelentkezo-kockazatokat/>



ÓBUDAI EGYETEM
BÁNKI DONÁT GÉPÉSZ ÉS
BIZTONSÁGTECHNIKAI MÉRNÖKI KAR

Köszönöm megtisztelő figyelmüket!

Óbudai Egyetem
Biztonságtudományi Doktori Iskola