

ISACA MÁSODIK SZERDAI ELŐADÁS

Varga Gábor

ChatGPT és más nyelvi modellek biztonsága az elvek és a gyakorlat szintjén

MI és innovációs szakértő

[linkedin.com/in/gaborvar](https://www.linkedin.com/in/gaborvar)

gaborvar@hotmail.com

2023. május 10.

Az ISACA Magyarországi Egyesület által szervezett Második szerdai előadásokon bármilyen eszközzel történő kép- vagy hangrögzítés tilos. (Idetartozik a mobiltelefonokkal készített kép- vagy hangfelvétel is). A jelen szabály be nem tartása szerzői- és szomszédos jogi jogsértés jogkövetkezményeit vonhatja maga után.



FELELŐSSÉG KIZÁRÁSA

Az elhangzott prezentációk tartalmát illetően az ISACA Magyarországi Egyesület nem vállal felelősséget az abban nyújtott információk aktualitásáért, helyességéért és teljességéért.

Az itt elhangzott információk nem feltétlenül egyeznek meg az ISACA Magyarországi Egyesület álláspontjával.

ChatGPT és más nyelvi modellek biztonsága az elvek és a gyakorlat szintjén

Varga Gábor

MI és innovációs szakértő

[linkedin.com/in/gaborvar](https://www.linkedin.com/in/gaborvar)

gaborvar@hotmail.com



*Néha zavar minket, s néha gépek zavarják,
Már zsebre vágjuk, s ha rosszul figyelsz, ő is zsebre vág.*

A ChatGPT jelenség

≡ Forbes

FORBES > INNOVATION > ENTERPRISE & CLOUD

What ChatGPT And Generative AI Mean For Your Business

CNN BUSINESS

Real estate agents say they can't imagine working without ChatGPT now

VentureBeat

Microsoft gives businesses a GPT boost in Teams and Viva Sales

CNN BUSINESS Markets Tech Media Success Perspectives Video

Microsoft is bringing ChatGPT technology to Word, Excel and Outlook

USA TODAY

New Bing with ChatGPT brings the power of AI to Microsoft's signature search engine

COMPUTERWORLD

UNITED STATES ▾

≡

NEWS

Microsoft's new Teams Premium tier integrates with OpenAI's GPT-3.5

Weeks after extending its multibillion dollar partnership with OpenAI, Microsoft has announced that new Teams AI capabilities will be underpinned by OpenAI's GPT-3.5 natural language model.

The Verge

Menu +

MICROSOFT / TECH / ARTIFICIAL INTELLIGENCE

Microsoft launches Azure OpenAI service with ChatGPT coming soon / ChatGPT is coming to this Azure service soon, as businesses get to use new AI models in their own apps.

The Verge

+ Follow

View Profile

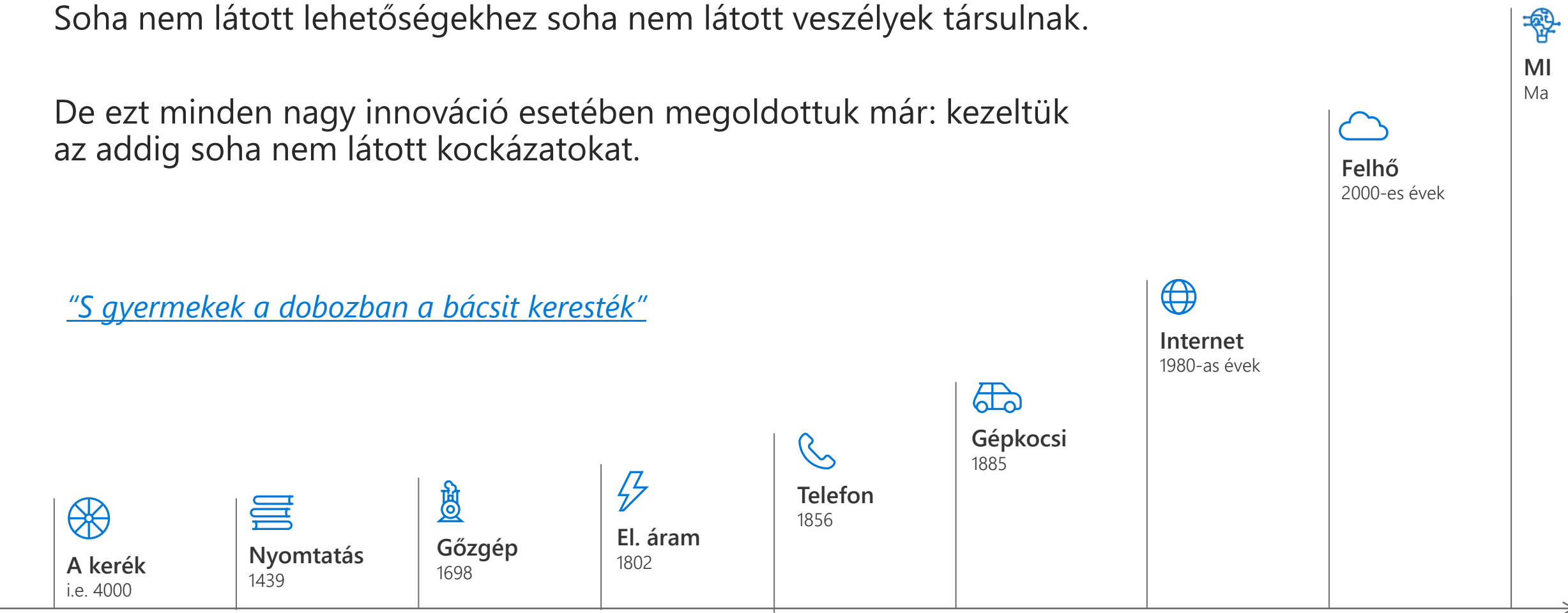
ChatGPT is now available in Microsoft's Azure OpenAI service

A mesterséges intelligencia a legnagyobb hatású innováció lehet az emberiség történetében

Soha nem látott lehetőségekhez soha nem látott veszélyek társulnak.

De ezt minden nagy innováció esetében megoldottuk már: kezeltük az addig soha nem látott kockázatokat.

"S gyermekek a dobozban a bácsit keresték"



Szabályozási analógia

Az automobil kitermelte a kockázatkezelés eszközeit:

- KRESZ
- oktatás társadalmi méretekben
- ittas vezetés tilalmának elfogadtatása a társadalommal
- töréstesztek és nyilvánosság
- műszaki vizsga, forgalmi engedély
- biztonsági öv – több évtizedes meggyőzés
- stb.

Ellenpélda: piros zászlós jelzőőr autó előtt

Hogyan alkalmazhatjuk ezeket a tapasztalatokat a mesterséges intelligencia szabályozásában?

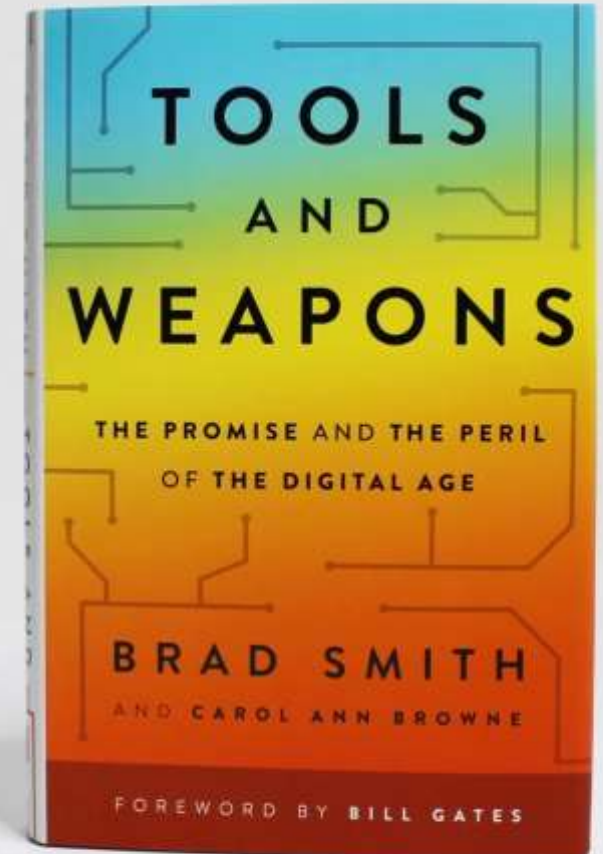


Eszköz vagy fegyver?

“Minél hatalmasabb az eszköz, annál nagyobb előnyökkel vagy károkkal járhat az alkalmazása...

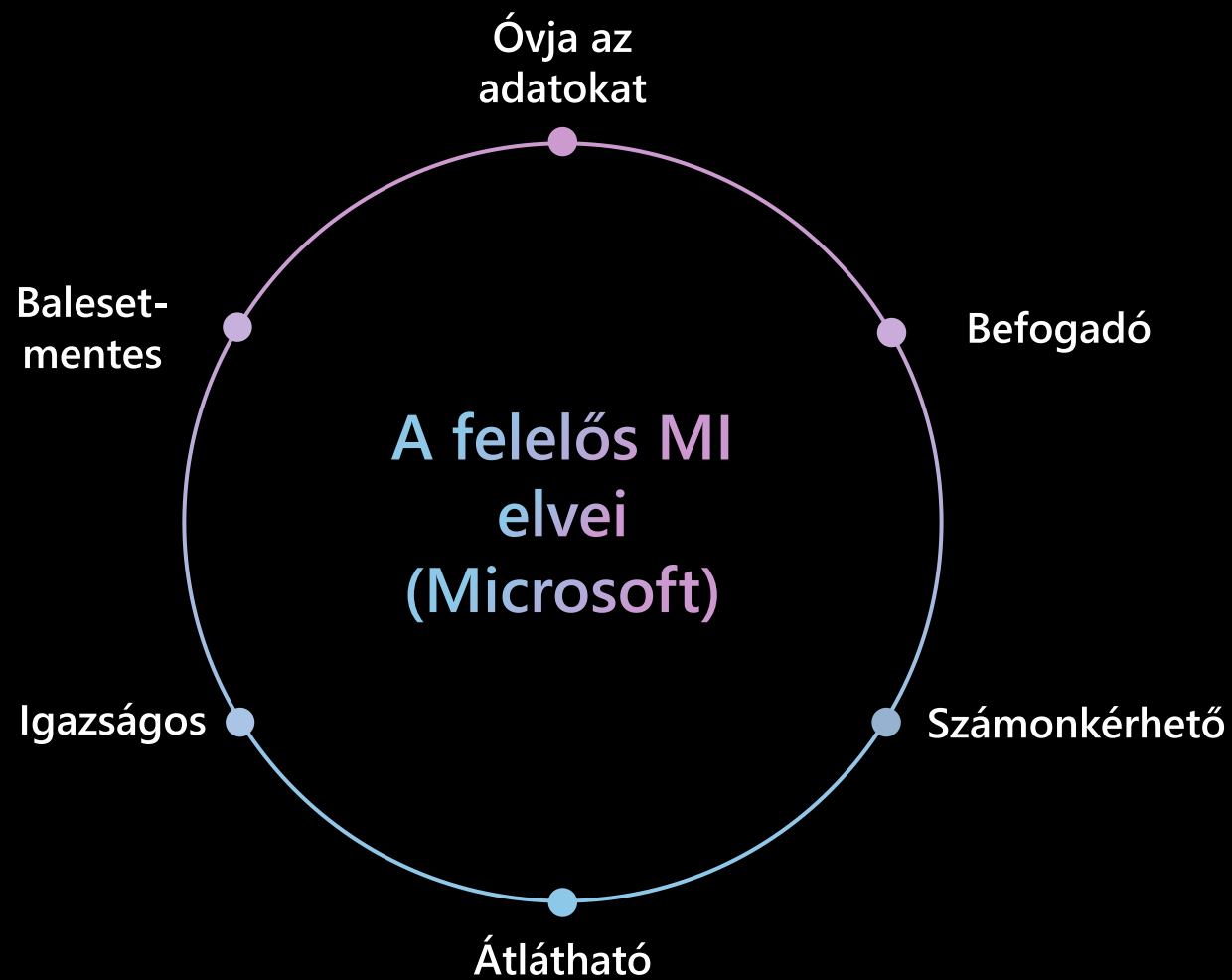
A technológiai fejlődés nem fog lelassulni. Annak a munkának kell felgyorsulnia, amellyel mederben tartható.”

Brad Smith
President and Chief Legal Officer, Microsoft



A mesterséges intelligencia felelős működése

"AI Alignment and AI Capabilities seem orthogonal. Progress in capabilities helps alignment. Alignment helps capability progress." - Sam Altman, Open AI



Az elvek érvényre jutásához:



Eszközök és folyamatok



Oktatás és gyakorlat



Szabályok, követelmények



Irányítás (Governance)

Gyakorlatias eszköz: Responsible AI Standard

<https://blogs.microsoft.com/wp-content/uploads/prod/sites/5/2022/06/Microsoft-Responsible-AI-Standard-v2-General-Requirements-3.pdf>

Accountability

A1: Impact Assessment
A2: Oversight of significant adverse impacts
A3: Fit for purpose
A4: Data governance and management
A5: Human oversight and control

Transparency

T1: System intelligibility for decision making
T2: Communication to stakeholders
T3: Disclosure of AI interaction

Fairness

F1: Quality of service
F2: Allocation of resources and opportunities
F3: Minimization of stereotyping, demeaning, and erasing outputs

Reliability & Safety

RS1: Reliability and safety guidance
RS2: Failures and remediations
RS3: Ongoing monitoring, feedback, and evaluation

Privacy & Security

PS1: Privacy Standard compliance
PS2: Security Policy compliance

Accessibility

I1: Accessibility Standards compliance

Példa a Responsible AI Standard alkalmazására

Goal A1: Impact assessment

Tools and Practices

Example Impact Assessment for Hospital Employee & Resource Optimization System (HEROS)

Stakeholders, potential benefits, and potential harms

The stakeholder analysis surfaces stakeholders and allows teams to consider how the system may impact them. Prompts related to the Goals of the Standard provide support for the identification of potential harms.

In the final section of the Impact Assessment, teams explore mitigations to the potential harms that have been identified.

[Microsoft's framework for building AI systems responsibly - Microsoft On the Issues](#)

Stakeholders, potential benefits, and potential harms

2.2 Identify the system's stakeholders for this intended use. Then, for each stakeholder, document the potential benefits and potential harms. For more information, including prompts, see the [guidance & activities deck](#).

Stakeholders	Potential system benefits	Potential system harms
Hospital Administrators (end user)	Automated shift scheduling reduces this stakeholder workload that would otherwise have to create schedules manually. This stakeholder benefits from a simplified task that involves making manual changes to the schedule if necessary and approving the automated schedules.	Reliability & safety: If this stakeholder finds that they must make a lot of manual changes because of repeated system errors, they may abandon the system all together. Accountability: if a lot of nurses have negative feedback this stakeholder may feel disproportionately accountable and spend a long-time reviewing schedules and making manual changes. Transparency & Accountability – If this stakeholder does not have enough information about nurses and their personal constraints/needs they may not be able to make informed changes and may be tempted to always approve the automatically generated schedules.
Nurses	Improves the general feeling of fairness. Nurses feel like shifts are allocated systematically and eliminates the fear of favouritism when compared to a human-operated system.	Fairness: Unsuitable shifts that do not account for personal circumstances. For example, a parent might not be able to work night shifts, or someone might need to work more hours for financial reasons. Reliability & safety: Our expert interviews revealed that nurses are more efficient when they work with the same team. Separation of medical teams could lead to unhappy, inefficient teams.

Mit tehetünk a MI védelmében

Készüljünk fel a támadások fajtáira és természetükre

[Failure Modes in Machine Learning](#) (Harvard, Microsoft)

Szándékos - nem szándékos; sértetlenség – bizalmasság – hozzáférhetőség, stb

Egyes támadási fajták kivédhetők a MI előtti módszerekkel

Kockázatelemzés

Hozzáférési jogosultságok

IT biztonsági higiénia – szoftverfrissítés, kétfaktoros autentikáció, stb

Mások MI specifikusak

Modellek, tanítási folyamat stb. támadása

[Best practices for AI security risk management - Microsoft Security Blog](#)
[AI-Security-Risk-Assessment/AI Risk Assessment v4.1.4.pdf at main · Azure/AI-Security-Risk-Assessment · GitHub](#)

Attack type	Likelihood	Impact	Exploitability
Extraction	High	Low	High
Evasion	High	Medium	High
Inference	Medium	Medium	Medium
Inversion	Medium	High	Medium
Poisoning	Low	High	Low

Eszközök

<https://github.com/microsoft/responsible-ai-toolbox>

Responsible AI Toolbox

Identify



Diagnose



Mitigate



Inform Actions



Error Analysis

Identify cohorts with high error rate versus benchmark and visualize how the error rate distributes



Model Interpretability

Interpret and debug model.



Unfairness Mitigation

Mitigate fairness issues



Exploratory Data Analysis

Understand dataset characteristics



Data Enhancements

Enhance your dataset and retrain model



Fairness Assessment

Assess model fairness by evaluating your model performance and predictions across different sensitive groups.



Counterfactual Analysis and What If

Generate diverse counterfactual explanations for debugging. Perform feature perturbations



Causal Inference

Understand the causal impact of your features on real-world outcomes



Counterfactual Analysis

Generate diverse counterfactual explanations for informing end users



Model
Comparison



Compare

Backward
Compatibility

A MI szunnyadó kérdéseket is felszínre hoz

Egyes kockázatok eddig is léteztek, de nem vettük észre őket. Kettős mérce: a gép előtt magasabbra tesszük a lécet, mert nagyobbbnak tartjuk az etikátlan döntések veszélyét

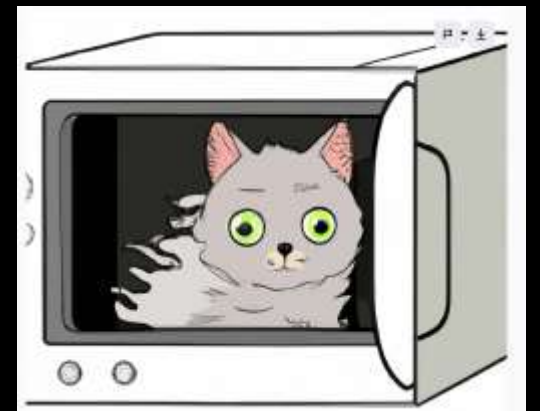
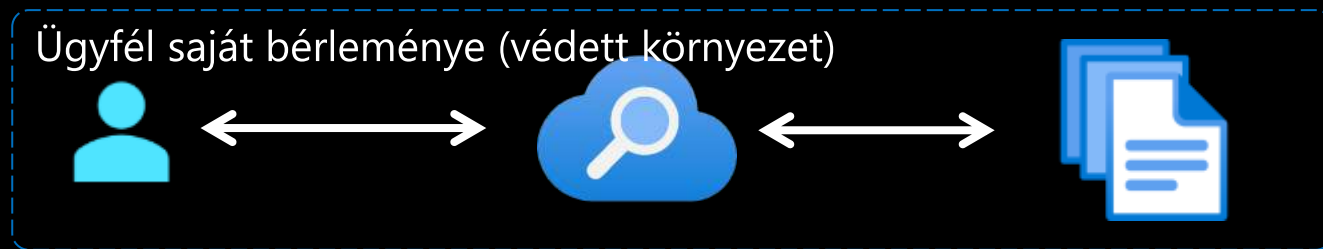
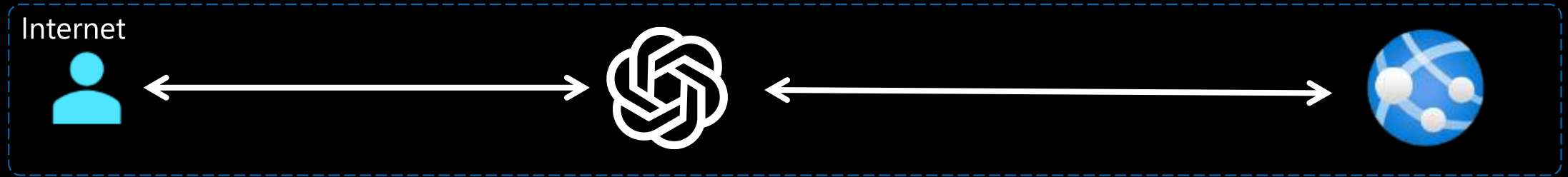
Three measurements of fairness: False positive rate, false negative rate, positive predictive value. Measure these 3 metrics for two subgroups. Each one is desirable. However cannot be satisfied simultaneously unless the base rate is the same. [You and AI – the challenges to making machines play fair](#)

Más kockázatok az új technikai lehetőségek miatt kapnak gyakorlati jelentőséget

[ChatGPT is a data privacy nightmare. If you've ever posted online, you ought to be concerned \(ampproject.org\)](#)
[\(16\) Nissenbaum's Theory of "Contextual Integrity" | LinkedIn](#)

Kockázati térkép – mit vállal és mit nem a szolgáltató?

Példa: OpenAI vs Azure OpenAI vs Azure Cognitive AI



Tudatos felelősségelosztás. Az ügyfélnek kell kezelnie egyes kockázatokat.

I Microsoft Azure Cloud

Nagyvállalati/intézményi bizalmi modell

A Te adatod a Tiéd

Az adatokat titkosítva tárolja az Azure előfizetésed

A továbbtanítások
(finomhangolások) adatait nem
használgák modellek
(elő)tanításához

Azure OpenAI szolgáltatás az ügyfél Azure előfizetésében
létesül
A továbbtanított modell az ügyfél Azure előfizetésében
marad és soha nem kerül át az alap modellekhez

Az adataidat átfogó
nagyvállalati/intézményi biztonsági
intézkedések védik, a szabályozói
megfelelés is cél

Ügyfél által kezelt kulcsok titkosítják (CMK)

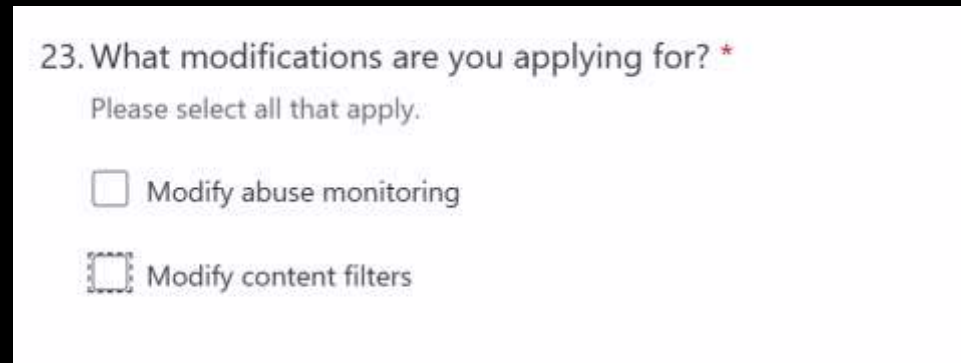
Magán VPN, szerepkör alapú hozzáférések (RBAC)

SOC2, ISO, HIPPA, CSA STAR Compliant

Tudatos döntések és hatásuk a felelős működésre

A tartalom szűrése bizonyos esetekben kikapcsolható.

Azure OpenAI Limited Access Review:
Modified Content Filters and Abuse Monitoring



23. What modifications are you applying for? *

Please select all that apply.

☐ Modify abuse monitoring

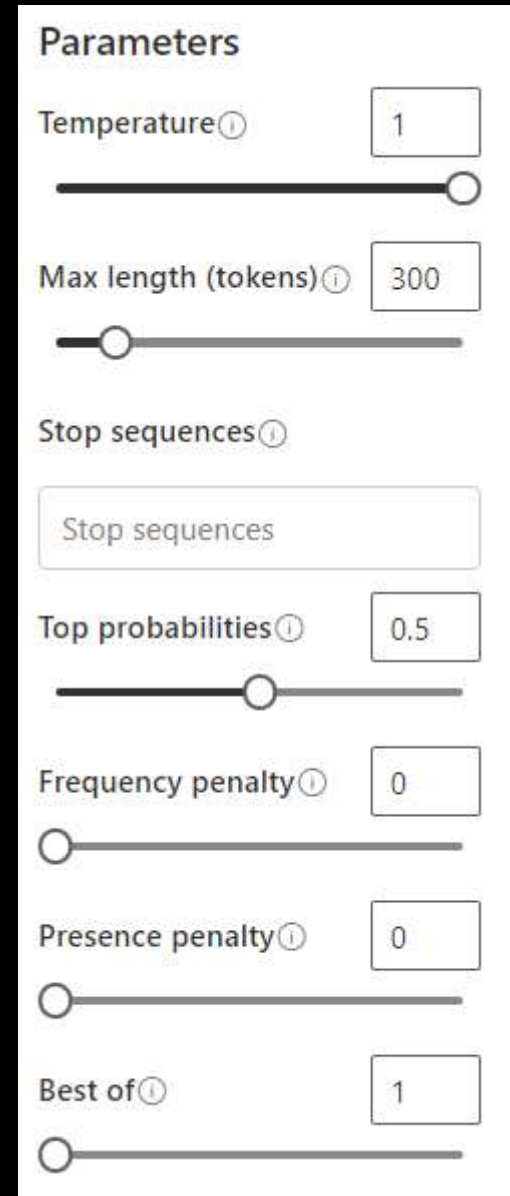
☒ Modify content filters

A válaszok “hőmérsékletének” megválasztása 0 és 1 között. Kreatív vagy kockázatkerülő legyen?

A felhasználó döntési lehetőségei és hatásuk a kockázatokra

- **Max_length** - The maximum number of tokens to generate in the completion. The token count of your prompt plus max_length can't exceed the model's context length.
- **Temperature** - What sampling temperature to use. Higher values means the model will take more risks. Try 0.9 for more creative applications, and 0 (argmax sampling) for ones with a well-defined answer.
- **Top_p** - An alternative to sampling with temperature, called nucleus sampling, where the model considers the results of the tokens with top_p probability mass. So 0.1 means only the tokens comprising the top 10% probability mass are considered

<https://learn.microsoft.com/en-us/azure/cognitive-services/openai/reference#completions>



The image shows a screenshot of the 'Parameters' section of the OpenAI API interface. It contains several settings with sliders and input boxes:

- Temperature**: A slider set to 1.0, with an input box showing '1'.
- Max length (tokens)**: A slider set to 300, with an input box showing '300'.
- Stop sequences**: A text input field containing 'Stop sequences'.
- Top probabilities**: A slider set to 0.5, with an input box showing '0.5'.
- Frequency penalty**: A slider set to 0, with an input box showing '0'.
- Presence penalty**: A slider set to 0, with an input box showing '0'.
- Best of**: A slider set to 1, with an input box showing '1'.



EU Ethics Guidelines for Trustworthy AI:

"a hátrányos megkülönböztetést el kell kerülni, mert sokszorosan negatív következményei lehetnek, a sérülékeny csoportok marginalizálódásától az előítéletesség és a kizárás erősödéséig."

A magyar nyelvet értő modellekre szükség van a világnyelveken nem beszélők esélyegyenlősége érdekében.

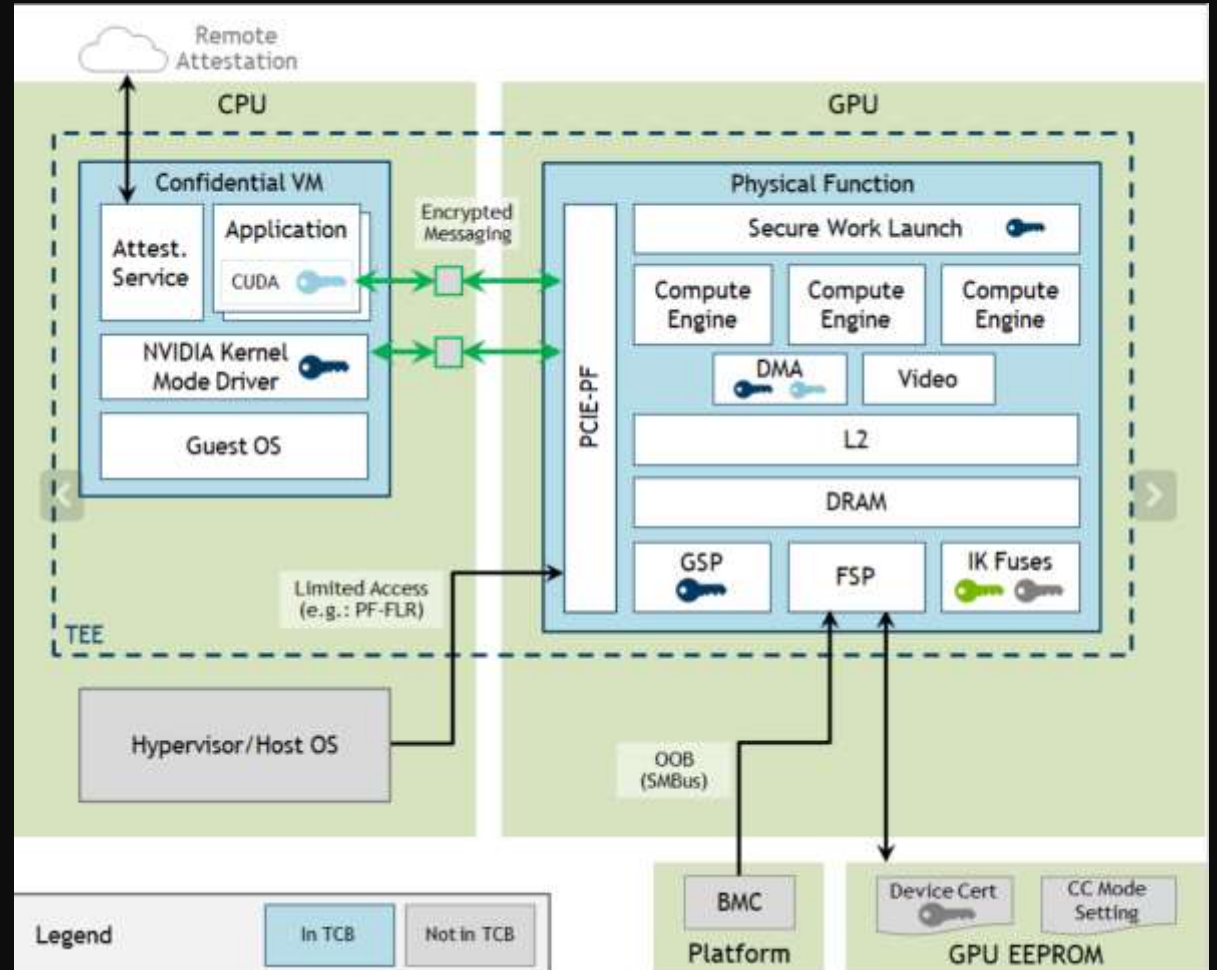
Bizalmi számítástechnika (Confidential Computing)

Többszereplős adatelemzés

Ha a szereplők nem engednek a többieknek hozzáférést a saját adataikhoz

Védje és tartsa titkosítva az adatokat, amelyek **használatban vannak**, miközben a gép számításokat végez rajtuk.

Nem kell megbíznod a szolgáltatóban. A bizalmi számítástechnika megakadályozza, hogy a rendszergazda hozzáférjen az adatokhoz. Erről kriptográfiai bizonyítékok felmutathatók.



Hogyan vehetünk részt az
ökoszisztémában?

A nagy nyelvi modellek ökoszisztémája

A MI másképp működik, mint a hagyományos IT

Extrém
infrastruktúra



Prompt
engineering



Továbbtanítás



Optimalizált
pipeline



Példa nagy nyelvi modellek többszereplős fejlesztésére

Értéklánc és ökoszisztéma létrehozásának módja

A megkülönböztetés eszközei

Promptok

“Barátságos, informatív
munkatárs vagy”

“Csak az adatokból
válaszolj”

“Ha nem találsz rá választ,
akkor válaszold ezt...”

Saját adatok

Belső tudásbázisok

Strukturált és
strukturálatlan források

Operációs és
tranzakciós adatok

GPT-3.5

GPT-4

ChatGPT

Azure OpenAI Service





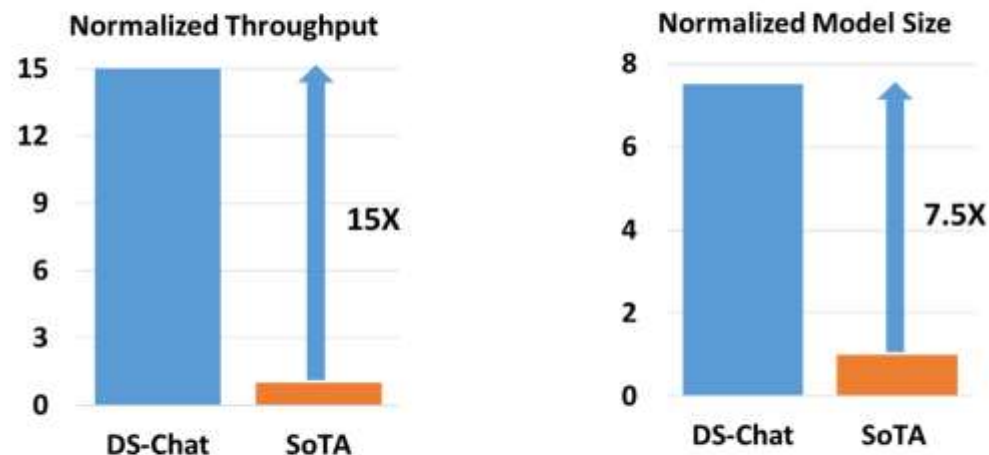
DEEPSPEED CHAT



Fast Training with Affordable Cost



Train 15X Faster and Scale to 5x Bigger Models than SOTA RLHFs



Easy-Breezy Training

A complete end-to-end RLHF training experience with a single click

High Performance System

Hybrid Engine achieves 15X training speedup over SOTA RLHF systems with unprecedented cost reduction at all scales

Accessible Large Model Support

Training ChatGPT-Style models with tens to hundreds of billions parameters on a single or multi-GPUs through ZeRO and LoRA

A Universal Acceleration Backend for RLHF

Support InstructGPT pipeline and large-model finetuning for various models and scenarios

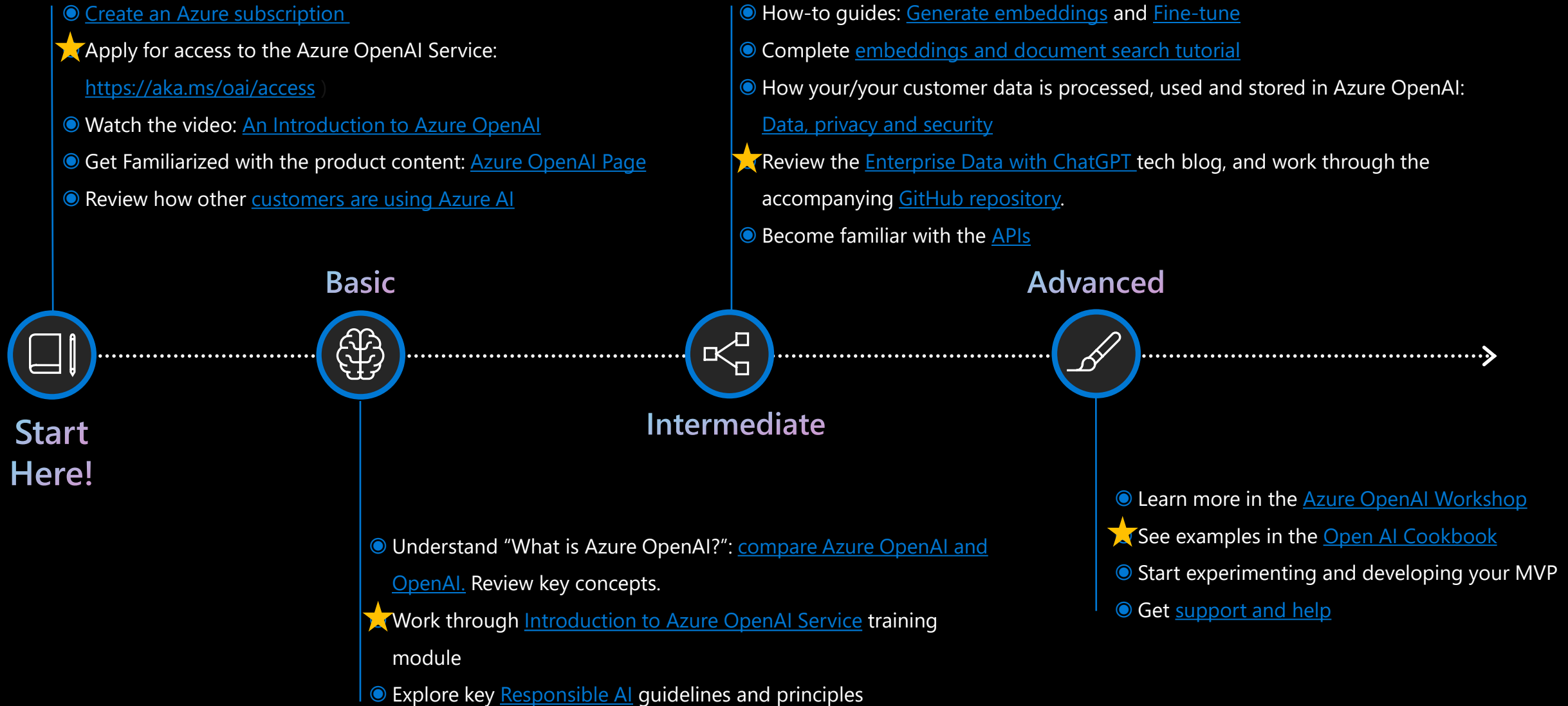
Kontextus

Context is King: In-context learning

Emergence

Egy modell, ami "szinte mindenre" képes. Alapjaiban változtatja meg a MI fejlesztésének dinamikáját, szerepköreit.

Azure OpenAI Service Learning Guide



„Ez a forradalom megállíthatatlan.

Ha megpróbálnánk, csak annyit

érnénk el, hogy a jövőt átengedjük

olyanoknak, akik bátrabban

szembe néznek a saját újításaik

következményeivel.”

„This revolution is unstoppable. Attempts to halt it would cede the future to that element of humanity more courageous in facing the implications of its own inventiveness.”

HENRY A. KISSINGER
ERIC SCHMIDT
DANIEL HUTTENLOCHER
on artificial intelligence
The Atlantic, August 2019



Összegzés

*Még kihasználatlan a védelem
eszköztára*

*Az etikus és felelős
használatról ne mondjunk le*

*Társadalmi mélységű
erőfeszítés*

*A hozadékok nagyok, indokolt
áldozni rá*

