

AI CYBERSECURITY SUMMIT

Mennyire vannak felkészülve a vállalatok az AI-agentek és robotok biztonságos használatára?

2026. 09. 08. 15:00

Innovation Hall

HUN · 30 perc

**Kerek István**

vezérigazgató · Everengine

**Benyovszky Máté**

alapító társelnök · Magyar Robotikai Szövetség

**Dr. Paniti Imre**

kutató, tudományos főmunkatárs · HUN-REN SZTAKI

**Béres Péter**

IT vezető · Sicontact Kft.

Dr. Paniti Imre, Benyovszky Máté, Béres Péter és Kerek István megvitatják az AI-ügynökök és robotika vállalati és fizikai környezetekbe történő integrálásának biztonsági kockázatait, gyakorlati kihívásait és irányítási követelményeit.

MIRŐL VOLT SZÓ

Dr. Paniti Imre hangsúlyozza a helyi célra készült szoftverek és harmadik féltől származó AI-ügynökök megkülönböztetésének szükségességét a adatáramlási kockázatok csökkentése érdekében. Míg sok szervezet nem rendelkezik egyértelmű stratégiával, néhányan „gyors nyereség” projekteket próbálnak megvalósítani anélkül, hogy alapvető adatirányítási vagy biztonsági politikákat hoznának létre. Például egy ügynök véletlenül törölheti a szerver tartalmát vagy az e-mail mappákat, ha nem megfelelően korlátozzák.

Béres Péter kiemeli, hogy az AI-ügynökök inkább emberi asszisztensekhez hasonlítanak, mint az Excelhez hasonló kiszámítható eszközökhöz, ezért a biztonsági gondolkodásnak az akciók és külső eszközök használatának megfigyelésére kell áttérnie. Kerek István az AI jelenlegi bevezetését a tűz történelmi uralmához hasonlítja, megjegyezve, hogy bár hatalmas hatalmat kínál, megköveteli a keretek kialakítását. Ezek az elemek azt sugallják, hogy a „biztonság tervezésből” keretrendszer nélküli gyors bevezetés jelentős sebezhetőségeket teremt mind a digitális, mind a fizikai térben.

KIEMELT MEGÁLLAPÍTÁSOK

- Béres Péter megjegyzi, hogy egy mesterséges intelligencia-ügynök úgy viselkedhet, mint egy emberi kolléga, ami szükségessé teszi a determinisztikus biztonságról a viselkedésalapú megfigyelésre való áttérést.
- Kerek István rámutat, hogy sok vállalat gyors eredményeket szeretne elérni adatstratégia nélkül, ami a digitális érettség lépéseinek kihagyásához vezet.
- Benyavszki Máté figyelmeztet, hogy a fizikai robotok, például a porszívók vagy a humanoidok, olyan kockázatokat jelentenek, ahol a szoftverhibák fizikai cselekvésekként nyilvánulhatnak meg.

ÜZLETI HATÁS

A prezentáció jelentős pénzügyi és működési kockázatokat azonosít, például az „LLM jacking” jelenséget, amikor ellopott munkamenet-tokenek engedély nélküli nagy értékű átutalásokhoz vagy hatalmas API-token költségekhez vezetnek. Ezek kivédésére a felszólalók homokozók használatát, a modelltesztelésre szánt dedikált gépeket és dinamikus elemzőeszközöket javasolnak a rosszindulatú URL-ek és IP-k felkutatására. A „VulnOps” és az ügynöki jogosultságok központosított kezelésének bevezetése praktikus módszerként szerepel az AI biztonságos skálázásához, miközben fenntartják a vállalati irányítást.

KÖVETKEZTETÉS

A felszólalók arra ösztönzik a hallgatóságot, hogy a biztonságot rejtett aggodalom helyett proaktív állapotba helyezték az AI-n belüli megoldások bevezetésével és a személyzet továbbképzésével. Kerek István azt javasolja, hogy Claude vagy Copilot modellekkel ismerkedjenek meg, hogy megértsék képességeiket, miközben szakértői segítséget kérnek az infrastruktúra terén. A végső üzenet az, hogy az AI-biztonságot folyamatos tanulási és alkalmazkodási folyamattal kezeljék, nem pedig egyszeri beállítással.

AI CYBERSECURITY SUMMIT

How Prepared Are Companies for the Safe Use of AI Agents and Robots?

2026. 09. 08. 15:00

Innovation Hall

HUN · 30 perc



István Kerek
CEO · Everengine



Benyovszky Máté
co-chair and founder · Magyar Robotikai Szövetség



Dr. Imre Paniti
Researcher, Senior Research Associate · HUN-REN SZTAKI



Péter Béres
IT leader · Sicontact Ltd.

Dr. Imre Paniti, Benyovszky Máté, Péter Béres, and István Kerek discuss the security risks, governance challenges, and practical implementation of AI agents in corporate and physical environments.

WHAT WAS DISCUSSED

Dr. Imre Paniti emphasizes the necessity of distinguishing between local purpose-built software and third-party AI agents to mitigate data leakage risks. He notes that while many organizations are eager to adopt AI, they often lack the maturity to manage the risks of uncontrolled agent behavior, such as an agent accidentally deleting server content or accessing unauthorized email accounts.

Péter Béres highlights that companies often pursue "quick win" projects without establishing a foundational data strategy or proper permission management. Benyovszky Máté expands the scope to physical robotics, warning that autonomous systems like drones or humanoid robots introduce physical safety risks alongside cyber threats. These elements collectively suggest that AI integration requires a "security by design" approach rather than reactive measures.

KEY TAKEAWAYS

- Benyovszky Máté shares an example where a robot vacuum's open port allowed remote control of all units of that model.
- Péter Béres mentions a case where an AI chatbot tricked a financial officer into transferring 63 million forints by mimicking a trusted colleague.
- Dr. Imre Paniti suggests using sandboxes and dedicated machines to isolate AI models from production environments.

BUSINESS IMPACT

The discussion identifies a significant gap between high private AI usage and low corporate adoption due to security fears. Organizations face economic risks from "LLM jacking," where attackers access API tokens to run up massive usage bills. To counter this, Péter Béres recommends implementing dynamic analysis tools to scan for malicious URLs and skills, while István Kerek advocates for upskilling employees to manage the risks of increased efficiency.

CONCLUSION

István Kerek concludes by urging the audience to move security from the "bottom drawer" to a central part of the workflow. He recommends that individuals and organizations engage with models like Claude or Copilot to learn about security and seek professional help to build robust guardrails.