

AI CYBERSECURITY SUMMIT

Securing applications against the rapidly evolving AI-fueled threat landscape

2026. 09. 08. 11:45

Innovation Hall

ENG · 30 perc

**Josh Goldfarb**

területi információbiztonsági vezető · F5, Inc.

**Tóth Zoltán**

F5 BDM · ALEF Distribution HU Kft

Josh Goldfarb, az F5, Inc. CISO-ja szerint a határterületi mesterséges intelligencia gyors fejlődése alapvetően felgyorsította a kibertámadások ütemét, ami a biztonsági alapelvekhez és a kockázatalapú irányításhoz való visszatérést igényli.

MI VOLT TÉMA

Josh Goldfarb elmagyarázza, hogy a mesterséges intelligencia lehetővé teszi a támadók számára a kifinomult szociális mérnöki tevékenységet, például többnyelvű, meggyőző adathalász e-mailek vagy vezetőket ábrázoló deepfake videók létrehozását. Továbbá a határvonalas mesterséges intelligencia lehetővé teszi a rosszindulatú szereplők számára, hogy gyorsabban azonosítsák a sebezhetőségeket és összekapcsolják a kiaknázásokat az alacsony súlyosságú problémákon keresztül, mint korábban lehetséges, például 2018-ban két évről 2026-ban mindössze négy órára csökkentve a kiaknázás kifejlesztéséhez szükséges időt.

A vita hangsúlyozza, hogy bár az AI új fenyegetéseket hoz létre, a behatolások többsége továbbra is emberi hibából, rossz beállításokból és árnyék AI komponensekből ered. Ezek ellen az intézményeknek túl kell lépniük a pusztán javításon, és átfogó védelmi stratégiát kell alkalmazniuk, amely magában foglalja az azonosítás kezelését, a hálózati szegmentálást és az AI-alapú észlelést. Ezek az elemek együtt azt jelentik, hogy a védekezőknek stratégiai módon kell használniuk az AI-t, hogy lépést tartsanak a modern automatizált támadások sebességével és kifinomultságával.

FONTOS TÉNYEK

- Az AI miatt az új biztonsági rés kihasználására rendelkezésre álló időtartam két évről négy órára csökkent.
- Az emberi hibák és a rossz beállítások továbbra is a legtöbb biztonsági incidens elsődleges kiváltói, annak ellenére, hogy egyre kifinomultabb AI eszközök jelennek meg.
- A biztonsági szabványoknak való megfelelés nem garantálja, hogy egy szervezet valóban védett legyen az új, aláírás nélküli támadásokkal szemben.

ÜZLETI HATÁS

A vállalati hatás az ellenőrizetlen AI jogi és pénzügyi kockázatainak kiemelésével érhető tetten, mint például az Air Canada esete, ahol egy chatbot hallucinációja kötelező érvényű politikává vált. A szervezeteknek vállalniuk kell az AI komponenseik kockázatát, prioritásként kezelve a legkritikusabb adatvagyon védelmét, és biztosítva, hogy a modellek pontos, biztonságos információkat szolgáltatassanak.

KÖVETKEZTETÉS

Josh Goldfarb azzal zárja, hogy arra ösztönzi a szervezeteket, hogy az AI-alapú fenyegetéseket integrálják meglévő biztonsági programjaikba, ahelyett, hogy pánikba esnének. Ajánl egy kockázatalapú megközelítést, amely a magas értékű eszközöket helyezi előtérbe, és folyamatosan értékeli az egész alkalmazás-verem állapotát, az AI-modell rétegtől egészen az alapvető infrastruktúráig.

AI CYBERSECURITY SUMMIT

Securing applications against the rapidly evolving AI-fueled threat landscape

2026. 09. 08. 11:45

Innovation Hall

ENG · 30 perc

**Josh Goldfarb**

Field CISO · F5, Inc.

**Zoltán Tóth**

F5 BDM · ALEF Distribution HU

Josh Goldfarb, Field CISO at F5, Inc., argues that the rapid advancement of frontier AI has fundamentally accelerated the timeline for cyberattacks, necessitating a return to core security hygiene and risk-based management.

WHAT WAS DISCUSSED

Josh Goldfarb explains that AI allows attackers to perform sophisticated social engineering, such as creating convincing phishing emails in multiple languages or deepfake videos of executives. Furthermore, frontier AI enables the rapid identification of vulnerabilities and the chaining of exploits across low-severity issues to compromise applications. For example, attackers can now develop exploits for newly discovered vulnerabilities in a matter of hours rather than years.

The discussion highlights that while AI increases attack sophistication, many breaches still stem from human error, misconfigurations, and shadow AI components. Josh Goldfarb emphasizes that organizations must move beyond simple patching to a defense-in-depth strategy involving identity management, network segmentation, and encryption. These elements combined mean that security must be proactive and risk-centric rather than reactive to individual vulnerabilities.

KEY TAKEAWAYS

- The time for attackers to exploit a new vulnerability has shrunk from an average of two years in 2018 to approximately four hours in 2026.
- Compliance with industry standards does not guarantee security, as organizations may remain vulnerable despite meeting specific regulatory checklists.
- Enterprises must own the risk of AI components, as seen in the Air Canada case where a chatbot's incorrect information became a binding legal policy.

BUSINESS IMPACT

The shift toward AI-powered attacks requires businesses to prioritize their most critical data and assets to manage risk effectively. Organizations must implement guardrails for AI models and the application code wrapping them to prevent data leakage or hallucinations. This approach helps mitigate the financial and legal ramifications of automated components making poor decisions or exposing backend systems.

CONCLUSION

Josh Goldfarb concludes by urging everyone, regardless of whether they are an AI company, to integrate AI-based threats into their existing security programs. He recommends starting with the most sensitive resources and applying security fundamentals to ensure the organization is not an easy target in a faster-moving threat landscape.