

Szabó Gergő¹ 

A generatív mesterséges intelligencia használata a kibertudatosítási képzésekben, 2. rész²

The Use of Generative Artificial Intelligence in Cyber Awareness Training, Part 2

Absztrakt

A tanulmány célja annak vizsgálata, hogy miként alkalmazható a generatív mesterséges intelligencia (GenMI) a kibertudatosítási képzésekben, különös tekintettel az oktatás személyre szabhatóságára, hatékonyságára és elfogadottságára. A kutatás elméleti része az ADDIE modell mentén tárgyalja a GenMI szerepét az oktatástervezésben, bemutatva annak lehetőségeit a tananyagfejlesztés, személyre szabás és visszacsatolás területén. Részletesen ismerteti a RAG-technológia szerepét is, különösen a tartalmi pontosság és az adatbiztonság aspektusaira fókuszálva. Az empirikus kutatás egy 109 fő részvételével végzett kérdőíves vizsgálatra épül, amelyet a Technology Acceptance Model (TAM) keretrendszere szerint szerkesztettem. A kérdőív többek között phishing szimulációkat, egy generált scénáriót, valamint attitűdkérdéseket tartalmazott.

Kulcsszavak: GenMI, oktatás, kibervédelem, RAG, TAM

This study examines how generative artificial intelligence (hereinafter: GenAI) can be applied in cybersecurity awareness training, with particular emphasis on personalisation, effectiveness, and user acceptance. The theoretical section follows the ADDIE instructional design model and explores the role of GenAI in educational planning, highlighting its potential in content development, adaptive learning, and feedback mechanisms. Special

¹ Nemzeti Közszerológati Egyetem, e-mail: sz.gergo628@gmail.com

² A Kulturális és Innovációs Minisztérium EKÖP-24-2-49 kódszámú egyetemi kutatói ösztöndíj programjának a Nemzeti Kutatási, Fejlesztési és Innovációs Alapból finanszírozott szakmai támogatásával készült.

attention is given to Retrieval-Augmented Generation (RAG) technologies, focusing on content accuracy and data protection considerations.

The empirical part of the research is based on a questionnaire survey conducted with 109 participants and designed according to the Technology Acceptance Model (TAM). The questionnaire included phishing simulations, an AI-generated cybersecurity scenario, and attitude-based questions. The aim of the study is to assess how non-expert users perceive and evaluate cybersecurity training materials generated by artificial intelligence, and to identify key factors influencing their acceptance and perceived usefulness.

Keywords: artificial intelligence, cybersecurity, awareness training, phishing

Bevezető

A kutatás első része elméleti és oktatástervezési szempontból közelítette meg a generatív mesterséges intelligencia szerepét a kiberbiztonsági képzésekben. Az elemzés fókuszában a tantervkészítés támogatása állt, különös tekintettel az ADDIE modell és a NIST CPLP keretrendszer egyes fázisaira. Bemutattam, hogy a GenMI képes hatékonyan támogatni az oktatási célú tartalomfejlesztést, ideértve nemcsak a szöveggenerálást, hanem az oktatási célokhoz igazított szcenáriók létrehozását is. Az érthetőség javítása érdekében a kutatási célkitűzéseket, kérdéseket és hipotéziseket a második rész elején mutatom be, ahol az empirikus kutatás részletes bemutatása is megtörténik.

A második részben az első részben felmerülő etikai és megbízhatósági kérdésekre felvetett megoldásként bemutatom a RAG-rendszerek működési elvét, valamint egy kérdőíves felmérés kvantitatív értékelésében bemutatom, hogyan értékelik a felhasználók a mesterséges intelligencia által előállított adathalász-szimulációkat és oktatási szcenáriókat. A kérdőív a Technology Acceptance Model (TAM) elméletére épült, és vizsgálta a tartalom érthetőségét, hasznosságát, a tanulási élményt, a technológia elfogadottságát, a jövőbeli használati szándékot és a válaszadók önértékelt tudásszintjének változását.

Kutatási célkitűzések és kutatási kérdések

A tanulmány célja kettős: egyrészt feltárni az elméleti háttér alapján, hogy a generatív mesterséges intelligencia miként támogathatja a személyre szabott, hatékony kibertudatosítási képzések kialakítását; másrészt empirikus úton vizsgálni, hogy a felhasználók miként viszonyulnak a GenMI által generált oktatási tartalmakhoz.

A kutatás során az alábbi célkitűzések mentén szerveződött az elemzés:

- KC1: A generatív mesterséges intelligenciával létrehozott kiberbiztonsági oktatási tartalmak elfogadottságának vizsgálata a felhasználók körében.
- KC2: A GenMI által generált tananyagok érthetőségének és tanulási élményének értékelése felhasználói visszajelzések alapján.
- KC3: Annak feltárása, hogy tapasztalható-e önértékelésen alapuló tudásszintváltozás egy GenMI-szcenárió elolvasását követően.

- KC4: Azoknak a tényezőknek az azonosítása, amelyeket a válaszadók a mesterségesintelligencia-alapú oktatás bevezetésének legnagyobb kihívásaiként jelölnek meg.
- KC5: Az MI-hez való hozzáállás különbségeinek feltérképezése demográfiai szempontok (például szakterület, életkor, iskolai végzettség) mentén.

A célkitűzésekhez kapcsolódón a tanulmány az alábbi kutatási kérdésekre keresi a választ:

- KK1: Milyen mértékben fogadják el a felhasználók a GenMI által létrehozott oktatási tartalmakat a kiberbiztonsági képzésekben?
- KK2: Hogyan értékeli a felhasználók a GenMI által generált tartalmak érthetőségét és tanulási élményét?
- KK3: Tapasztalható-e tudásszintbeli változás egy GenMI által generált kiberbiztonsági szcenárió elolvasását követően?
- KK4: Milyen tényezőket tartanak a válaszadók a legnagyobb kihívásnak a mesterségesintelligencia-alapú oktatás bevezetésében?
- KK5: Vannak-e különbségek az MI-hez való hozzáállásban a válaszadók szakterülete, életkora vagy iskolai végzettsége szerint?

Hipotézisek

A fenti kutatási kérdésekhez az alábbi hipotézisek tartoznak:

- H1: A válaszadók többsége elfogadja a GenMI által generált kiberbiztonsági oktatási tartalmakat hiteles és hasznos tananyagként.
- H2: A résztvevők a GenMI által generált tananyagot érthetőnek és tanulásra ösztönzőnek találják.
- H3: A GenMI szcenárió elolvasását követően a válaszadók magasabbra értékeli saját képességüket az adathalászat felismerésében.
- H4: A leggyakrabban említett kihívás a GenMI-alapú oktatásban a generált tartalom megbízhatósága és az oktatási rendszerbe vetett bizalom hiánya.
- H5: A GenMI-hez való pozitív hozzáállás szignifikánsan magasabb az informatikai szakterületen dolgozók körében, mint más területeken.

A hipotézisek tesztelése kvantitatív adatfeldolgozással történik, a kérdőíves válaszok statisztikai elemzése alapján. Az egyes hipotézisek alátámasztottságát a harmadik fejezet részletesen bemutatja.

Módszertan

A tanulmány két szinten közelíti meg a kutatási kérdéseket: egyrészt elméleti, másrészt empirikus alapon. Az elméleti rész célja a generatív mesterséges intelligencia oktatásban betöltött szerepének bemutatása, különös tekintettel a személyre szabható, adaptív tanulási folyamatok lehetőségeire a kiberbiztonsági tudatosság fejlesztésében. Ennek érdekében szakirodalmi áttekintést végeztem a releváns nemzetközi és hazai

források alapján, külön figyelmet fordítva a GenMI és a kiberbiztonsági oktatás kapcsolódási pontjaira.

Az empirikus kutatás online kérdőív formájában valósult meg, amelyet a Technológiaelfogadási Modell (Technology Acceptance Model, TAM) mentén építettem fel. A kérdőív célja az volt, hogy valóságghű, generatív MI által készített adathalászeszteken és komplexebb kiberbiztonsági szcenáriókon keresztül vizsgálja a felhasználói élményt, az észlelt hasznosságot, az érthetőséget, valamint a bizalmat és az elfogadottságot. A kitöltés során lehetőség nyílt szöveges válaszadásra is, amely az élmény minőségi értékelését segítette.

A kvantitatív adatok feldolgozása a leíró statisztikákon túl keresztábra-elemzéssel és Khi-négyzet-próbával történt, hogy feltárható legyen, mutatkozik-e szignifikáns különbség a különböző demográfiai csoportok (például nem, szakterület) között az MI-vezérelt oktatás elfogadottságában. Külön vizsgálat készült az IT-területen dolgozó és a nem informatikai területről érkező válaszadók válaszainak összevetésére is.

Információ-visszakereséssel kiegészített szöveggenerálás

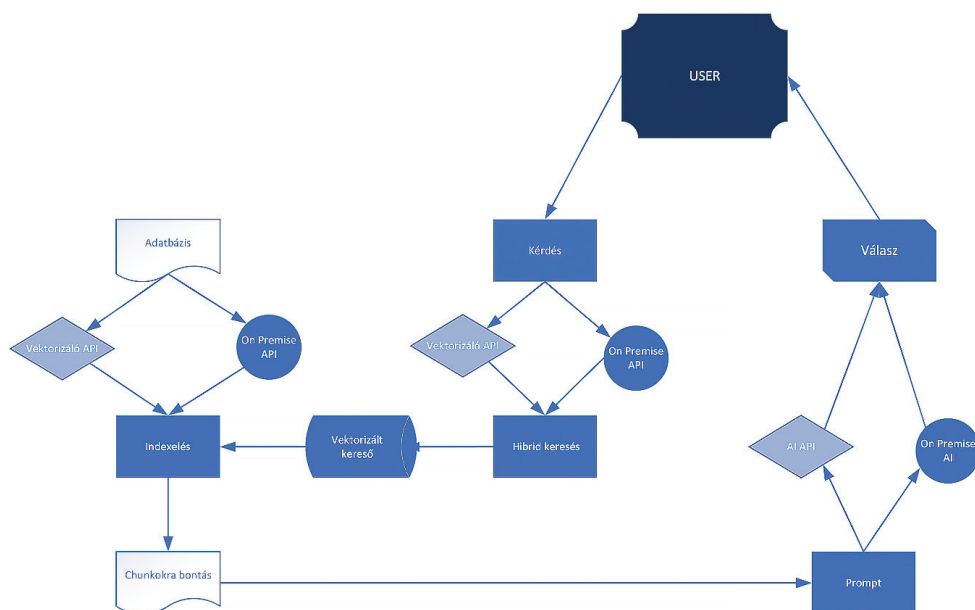
A RAG-rendszerek lényege, hogy a mesterséges intelligencia nem a tanított tudásból dolgozik, hanem valós időben keres információkat egy belső tudásbázisból. A vállalati képzések GenMI segítségével történő támogatásánál választható felhőalapú, szolgáltatott szoftver (*software as a service*, SaaS) vagy házon belüli (*on-premise*) RAG-megoldásrendszer. Különösen a szigorúan szabályozott iparágakban működő vállalatok preferálják, hogy az érzékeny adatokat saját infrastruktúrában tartsák, ahelyett, hogy azokat külső szolgáltatóra bíznák. Bár a legtöbb kész *awareness* tréning platform elsősorban SaaS-formában érhető el, bizonyos szolgáltatók kínálhatnak privát felhő- vagy virtuális *appliance* telepítést nagyvállalati ügyfeleknek egyedi megállapodás keretében.³

Alternatív megközelítés, hogy a vállalat saját maga épít RAG-alapú chatbotot vagy oktatórendszert nyílt forráskódú eszközökkel. A LangChain keretrendszer például lehetővé teszi egyedi, belső tudásbázison nyugvó chatbot kifejlesztését, amely házon belüli környezetben futó NNYM-ekkel működik.⁴ Ilyen *self-hosted* megoldásnál a cég egy nyílt nyelvi modellt, például a Llama 2-t, finomhangolhat saját adatokkal, vagy *retrieval-augmented generation* megközelítést alkalmazhat; az MI valós időben keres ki információkat egy belső tudásbázisból, ahelyett hogy minden vállalati adatot betanítanak a modellre.⁵ Ez a megközelítés hibrid modellként működik, mivel a generatív modell lehet felhőben, használhatja például az OpenAI API-t, de maga az adatkeresés és -tárolás helyben történik. Az 1. ábra szemlélteti a RAG általános működési elvét.

³ TitanHQ 2025.

⁴ HAYDOCK 2024.

⁵ HAYDOCK 2024.



1. ábra: A RAG általános működése

Forrás: NKE ÁNTK Kiberbiztonsági mesterképzés Hírszerzés a kibertérben tantárgy, Vadász Pál egyetemi előadása alapján (2025. március 1.)

A *hybrid* vagy *on-premise* megoldások előnye, hogy a vállalati dokumentumokat, tudásbázist helyben, például egy vektortárban kezelik, míg a generatív modell kizárólag lekérdezés útján fér hozzá az információhoz, szükség esetén. Ezáltal az üzleti titok nem kerül tartósan a modell paramétereire közé, csupán a kérdés megválaszolásának idejére.⁶ Ezzel elkerülhető egy problémás helyzet, mert ha érzékeny adatokkal tanítják be a modellt, később nem garantálható, hogy a betanításra használt adatok nem bukkannak-e fel valakinek a beszélgetésében, ami a GDPR szerint jogellenes adatkezelésnek minősülhet megfelelő törlési mechanizmus hiányában.⁷ Számos vállalat dönt a nyílt forráskódú nagy nyelvi modellek házon belüli alkalmazása mellett, ilyen esetben a modell betanítása és futtatása teljes mértékben az ő környezetükben történik, bár ennek erőforrásigénye jelentős lehet.

A RAG-rendszereknek biztosítaniuk kell az adatlekérdezések naplózását, és azt, hogy szükség esetén pontosan visszakereshető legyen, mely forrás alapján készült a generált válasz. Ez egyrészt a megbízhatóság miatt fontos, például azonosítható legyen, hogy a chatbot mely belső szabályzat melyik pontja alapján javasolt választ, ami nemcsak a megbízhatóságot, hanem az auditálhatóságot is támogatja. Sok megoldás ezért hivatkozásokat vagy forráslinkeket is mellékel az MI által generált válaszokhoz. Technikai szempontból elengedhetetlen, hogy a rendszer képes legyen az adatok osztályozására, érzékeny adatot ne jelenítsen meg illetéktelen felhasználónak még

⁶ HAYDOCK 2024.

⁷ HAYDOCK 2024.

akkor sem, ha az a tudásbázisban szerepel. Az AI Act várhatóan olyan előírásokat hoz, hogy a generatív modellek outputjainak bizonyos kockázatait, mint például a hallucináció, téves információ kezelve legyenek,⁸ erre a RAG részben válasz, hiszen a valós idejű információkeresés csökkenti a modellek hallucinációját azzal, hogy tényleges dokumentumokra támaszkodik.⁹ Mindazonáltal a vállalatoknak óvintézkedéseket kell bevezetniük, például felhasználói promptszűrés, nehogy üzleti titkot kérdezzenek be figyelmen kívül nyilvános modellhez, illetve hozzáféréskontroll az MI-eszközök felett.

Összességében elmondható, hogy a korszerű SaaS-plattformok vállalati kiadásai adatvédelmi szempontból egyre jobban kontrolláltak, szerződésben rögzítik az adatok tulajdonjogát, a modell nem tanul tovább az ügyfél adataiból, és gyakran tanúsítottak (SOC 2,¹⁰ ISO27001¹¹). További adatvédelmi garanciát jelenthet, ha a vállalat saját üzemeltetésű megoldás mellett dönt, ami teljes adatkontrollt ad, cserébe magasabb költséggel és komplexitással.

A RAG-alapú oktatórendszerek elsődleges célcsoportja a céges dolgozók széles köre, az újonnan csatlakozó munkavállalóktól kezdve a nem informatikai területen dolgozó alkalmazottakig. A platformok személyre szabott és interaktív tartalommal igyekeznek mindenki számára érthetővé és relevánssá tenni a kibervédelmi ismereteket. Generatív MI segítségével például valóság-hű támadási szimulációk készülnek: e-mailek, hanghívások, SMS-ek, QR-kódos támadások stb. A Keepnet Labs platformja MI-vel generált *phishing*, *vishing*, *smishing* és *quishing* (QR-kódon keresztül történő adathalászat) szimulációkat kínál, amelyeket az adott dolgozó munkakörére és korábbi teljesítményére szabnak.¹²

Hasonlóan a Hoxhunt nevű tréningplatform MI-vezérelt, adaptív mikroképzéseket nyújt, a rendszer minden egyes felhasználó viselkedése alapján egyénileg időzített és testre szabott adathalászat-e-maileket küld, ezzel kutatások szerint akár negyvenszeres növekedést is elér a felhasználói aktivitásban, és hosszú távú viselkedésváltozást ér el a hagyományos módszerekhez képest.¹³ Az MI a felhasználók eredményeit elemelve képes gyenge pontokat azonosítani és ezekre fókuszáló tartalmat ajánlani. Például ha egy osztályon sokan dőlnek be egy bizonyos típusú adathalásznak, a rendszer több ilyen példát küld, majd azonnali visszajelzést ad a tanulóknak a hibáikról, segítve a folyamatos fejlődést.¹⁴

Ezek a megoldások új belépők számára alapozótréninget nyújtanak, például jelzőkezelés, e-mail-biztonsági alapok, míg a régóta ott dolgozók számára friss, aktuális fenyegetési példákat mutatnak be, fenntartva a figyelmet, és elkerülik az úgynevezett tanulási kimerültséget vagy tréningfáradtságot, amely gyakran a túl sűrű vagy monoton oktatási programok esetén jelentkezik.

⁸ EU Artificial Intelligence Act. The AI Act Explorer 2025.

⁹ InterSystems 2025.

¹⁰ Az AICPA (American Institute of Certified Public Accountants) által meghatározott szabvány, amely azt értékeli, hogy a szolgáltató hogyan kezeli az ügyféladatokat az adatbiztonság, adatkezelés, rendelkezésre állás, titkosság és integritás szempontjából.

¹¹ Nemzetközi szabvány az információbiztonság-irányítási rendszerekre (ISMS), amely biztosítja, hogy a szervezet rendszerszinten kezelje és védje az információkat a jogosulatlan hozzáférés, módosítás vagy elvesztés ellen.

¹² Keepnet 2024.

¹³ TitanHQ 2025.

¹⁴ Keepnet 2024.

Egyes RAG-rendszereket kiterjesztettek a haladó felhasználók igényeire is. Bár a hagyományos biztonságtudatossági platformok (KnowBe4, Hoxhunt stb.) főként az általános munkaerőt célozzák, a közelmúltban olyan generatív MI-alapú eszközök váltak elérhetővé, amelyek közvetlenül a SOC-csapat munkáját támogatják. A 2023-ban bemutatott Microsoft Security Copilot jó példa erre: a GPT-4 alapú asszisztens a Microsoft globális fenyegetésfelderítési adatain (napi 65 billió biztonsági jelzés) és a vállalat saját biztonsági modelljén keresztül ad valós idejű biztonsági eseményelemzéseket.¹⁵

A Security Copilot képes bonyolult támadási mintázatokat természetes nyelven összefoglalni, javaslatokat adni az incidensek kivizsgálásához, és integrálódik a meglévő védelmi eszközökhöz, például a Microsoft Defender vagy Sentinel SIEM.¹⁶ A security operations center (SOC) támogatására tervezett funkciók közé tartozik, hogy ilyen MI-asszisztenssel a biztonsági elemzők rövid szöveges utasítások formájában, például csevegőfelületen keresztül kérhetnek le fenyegetettségi információkat, log elemzéseket, vagy akár komplett incidensriportokat. A generatív modell a háttérben lekérdez több adatforrást, például vállalati naplókat, ismert sebezhetőségi adatbázisokat, ami támogatja a gyorsabb döntéshozatalt és csökkenti az incidensek kezelési idejét.¹⁷

Friss kutatási eredmények is alátámasztják e megközelítés hasznát: az *IntellBot* nevű kísérleti RAG-chatbot például számos adatforrásból gyűjt össze kibertámadásokkal kapcsolatos fenyegetéseket, ismert sérülékenységeket, aktuális incidenseket, és a felhasználók kérdéseire az aktuális tudásbázisból generál választ. Ezzel azonnali hozzáférést nyújt releváns információkhoz, segítve a fenyegetéselemzést és az incidenskezelést, időt spórolva meg a szakembereknek.¹⁸

A SOC-támogató generatív MI-rendszerek jellemzően nem sorolhatók a klasszikus értelemben vett oktatási rendszerek közé, hanem inkább operatív eszközök. Épp ezért a SOC-orientált funkciókat az oktatási rendszerek dokumentációjában rendszerint külön jelölik. Ugyanakkor oktatási hozadékuk is van, támogatják a kevésbé tapasztalt elemzők tudásának bővítését vagy kompetenciafejlesztését azáltal, hogy javaslatokat és magyarázatokat fűznek a talált eredményekhez.¹⁹ Például a CrowdStrike Charlotte MI asszisztense a vállalat végponti védelmi platformjába épül be, és automatikusan képes riasztásokat kezelni és reagálni azokra (*triage*). A gyártó állítása szerint 98% feletti pontossággal leválogatja a triviális vagy ismert mintázatú riasztásokat, így a SOC-csapat csak a valóban kritikus esetekre fókuszálhat, ezzel átlagosan heti 40 órányi manuális munkát spórol meg a csapatnak.²⁰

A Microsoft Security Copilothoz hasonlóan a Charlotte MI is generatív interfészt kínál, például strukturált incidens-összefoglalót generál, amely támogatja az utólagos elemzést és a tudásmegosztást.²¹

¹⁵ JAKKAL 2023.

¹⁶ CHANDROTH 2025.

¹⁷ ARIKKAT et al. 2024.

¹⁸ ARIKKAT et al. 2024.

¹⁹ REED–RAY 2024.

²⁰ CrowdStrike 2025.

²¹ COLUMBUS 2023.

Összegzőként megállapítható, hogy a RAG-megközelítést alkalmazó generatív MI-rendszerek nemcsak az általános munkavállalók biztonságtudatosítási képését, hanem a kiberbiztonsági szakemberek napi munkáját és folyamatos szakmai fejlődését is képesek támogatni. Az ilyen eszközök bevezetése hozzájárulhat a gyorsabb incidenskezeléshez, a munkaterhelés csökkentéséhez, valamint a szervezeti tudás hatékonyabb hasznosításához.

A generatív mesterséges intelligencia által készített feladatok elfogadottsága Magyarországon

A technológiai megoldások elfogadottsága és hatékonysága nem csupán elvi kérdés, hanem mérhető élmény is. A generatív mesterséges intelligencia oktatási célú alkalmazása különösen érzékeny terület, mivel egyszerre érinti a tanulási folyamat minőségét, a bizalom kérdését és a felhasználói önreflexiót. A kérdőív válaszai nem pusztán véleményeket tükröznek, hanem különböző demográfiai és hozzáállásbeli mintázatokat is kirajzolnak. E fejezet célja ezen mintázatok feltérképezése és annak bemutatása, hogy a GenMI által előállított tartalmak milyen mértékben váltottak ki értelmezhető reakciókat a válaszadók körében. A statisztikai elemzés nemcsak az elfogadottságot, hanem a tanulási élményt, a viselkedésbeli változást és az MI-hez való viszonyulás sokféleségét is megvilágítja.

A kérdőív felépítése és célja

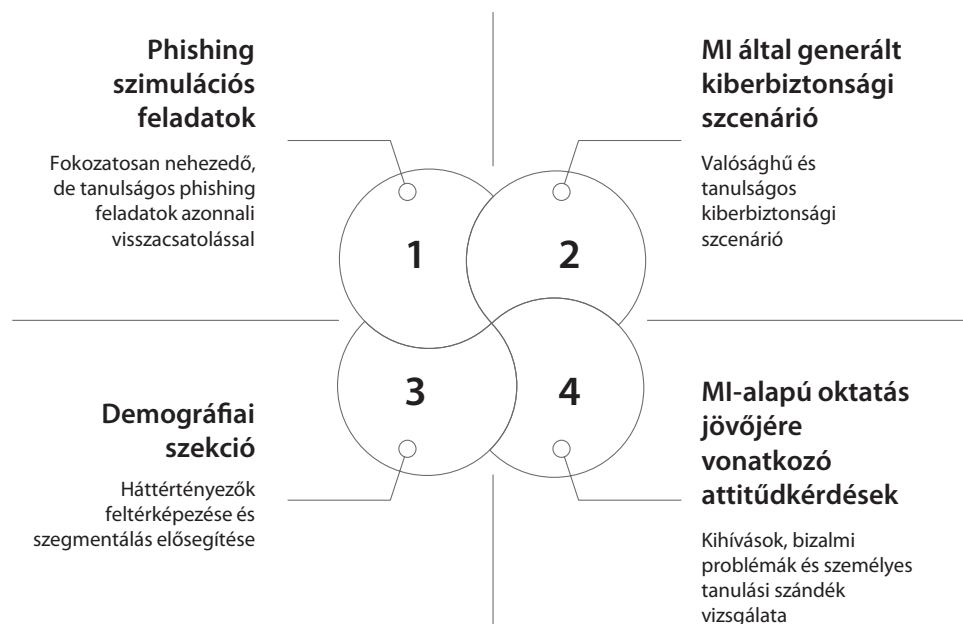
A kérdőíves kutatás célja az volt, hogy feltérképezze, miként viszonyulnak a válaszadók a generatív mesterséges intelligencia által létrehozott oktatási tartalmakhoz a kibertudatosítási képzések kontextusában. A kérdések kialakítását a Technology Acceptance Model elméleti kerete inspirálta,²² amely szerint egy új technológia elfogadottsága elsősorban két pszichológiai tényezőtől függ: az észlelt hasznosságtól és az észlelt használati könnyűségtől. Előbbi azt tükrözi, hogy a felhasználó mennyire látja értelmesnek és eredményesnek az adott technológia alkalmazását, míg utóbbi arra utal, mennyire tartja azt intuitívnak, érthetőnek vagy gyorsan elsajátíthatónak.

A kérdőívben az észlelt hasznosság két konkrét szempont mentén jelent meg, a tanulási érték és a jövőbeli alkalmazhatóság alapján. Ezek azért fontosak, mert nemcsak a GenMI-alapú szcenárió azonnali megítélését tükrözik, hanem azt is, hogy a válaszadók hosszabb távon mennyire látják integrálhatónak ezt a technológiát a kiberbiztonsági képzésekbe. Az érthetőségi faktor vizsgálata, a szcenáriók világossága és magyarázatai alapján, az észlelt könnyű használatot közelítette meg. Ezen túlmenően kérdések irányultak az MI-alapú oktatás elfogadási szándékára és a válaszadók érzelmi reakcióira is, így a kérdőív az elfogadás és érdeklődés, bevonódás összetett pszichológiai képét igyekezett feltárni.

²² DAVIS 1989.

A kérdőív négy nagyobb egységből épült fel:

- Három egymásra épülő *phishing* szimulációs feladat, amelyek fokozatosan komplexebbek lettek, de nem haladták meg a kitöltők általános információs műveltségét. A feladatok zárt végűek voltak, 4 válaszlehetőséggel. A válaszadók a visszacsatolások után három zárt kérdésre válaszoltak egy 1–5 fokozatú Likert-skálán (1 = egyáltalán nem, 5 = teljes mértékben), amelyek a feladat érthetőségére, a magyarázat világosságára és a tanulási élmény mértékére irányultak. A kérdések célja nem a válaszok helyességének vizsgálata volt, hanem a generált tartalom pedagógiai értékének mérése.
- Ezt követte egy MI által generált, interaktív kiberbiztonsági szcenárió, amely egy közel valósághű helyzetet modellezett. Az esettanulmány elkészítésekor igyekeztem egyensúlyozni a részletességet és az elolvasáshoz szükséges időt, aminek kihatása van arra, hogy mennyire lesz valósághű az esettanulmány. Ehhez kapcsolódva a válaszadók a szcenárió hasznosságát, érthetőségét, tanulságosságát, valamint saját érzéseiket és észrevételeiket is megoszthatták a fenti Likert-skálás módszerrel.
- A következő szakaszban az MI-alapú oktatás jövőjére vonatkozó attitűdkérdések lettek feltéve, amelyek a legnagyobb kihívásokra, bizalmi problémákra és személyes tanulási szándékokra kérdeztek rá.
- Zárásként egy demográfiai szekcióból, amely tartalmazta az életkort, nemet, legmagasabb iskolai végzettséget, gazdasági aktivitást, lakóhelytípust és a szakmai háttér megadását.



2. ábra: A kérdőív felépítése

Forrás: a szerző szerkesztése a Napkin.ai segítségével

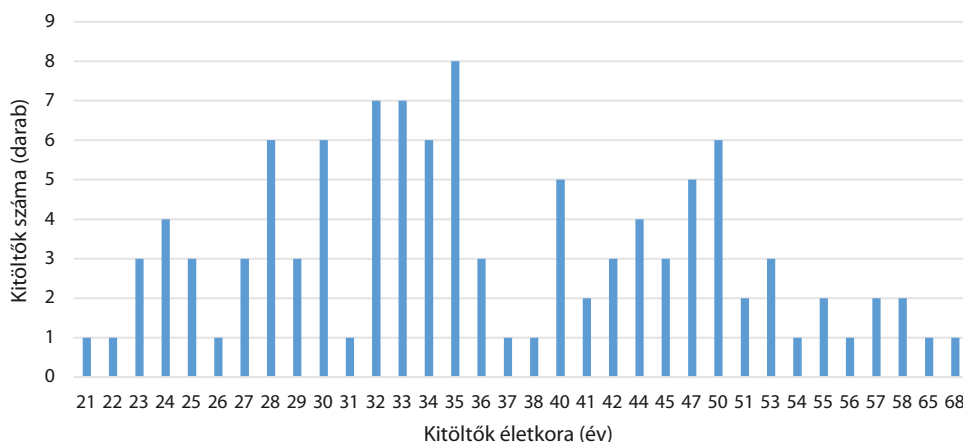
A kérdőívet Facebookon osztottam meg, kitölteni 2025. március 10. és április 10. között lehetett, ez idő alatt 109 fő töltötte ki, nem reprezentatív, de sokszínű minta alapján. A válaszadók között szerepeltek informatikai és kiberbiztonsági háttérrel rendelkezők is (27 fő), azonban a kérdőív célja nem szakmai mélységű értékelés, hanem a generált tartalmak általános használhatóságának és érthetőségének vizsgálata volt. A válaszok elemzése ennek megfelelően minden résztvevő válaszait figyelembe veszi, de az értelmezés során a nem szakértői nézőpontot tekintettük elsődlegesnek.

A válaszadók demográfiai jellemzői

A kérdőívet összesen 109 fő töltötte ki, önkéntes alapon. A minta nem reprezentatív, azonban célja nem is statisztikai általánosíthatóság, hanem annak vizsgálata volt, hogy a különböző háttérrel rendelkező felhasználók hogyan viszonyulnak a generatív mesterséges intelligencia által létrehozott kiberbiztonsági tartalmakhoz. A minta összetétele változatos, így alkalmas az elfogadottság, a tanulási élmény és az attitűdök közötti mintázatok feltérképezésére.

A válaszadók életkora 21 és 68 év között szóródott, az átlagéletkor 37,7 év, a szórás 10,5 év volt (3. ábra). A minta tehát alapvetően aktív korú felnőtteket reprezentál, akik potenciális célcsoportjai lehetnek egy valós vállalati vagy intézményi kibertudatosító programnak.

Egy válaszadó esetében az életkor mező nem numerikus formátumban szerepelt („3 x 5”), ezért ezt az adatpontot kizártam az életkorra vonatkozó statisztikai számításokból. Az érintett kitöltő többi válasza értelmezhető és konzisztens volt, így azok az elemzés részét képezték.



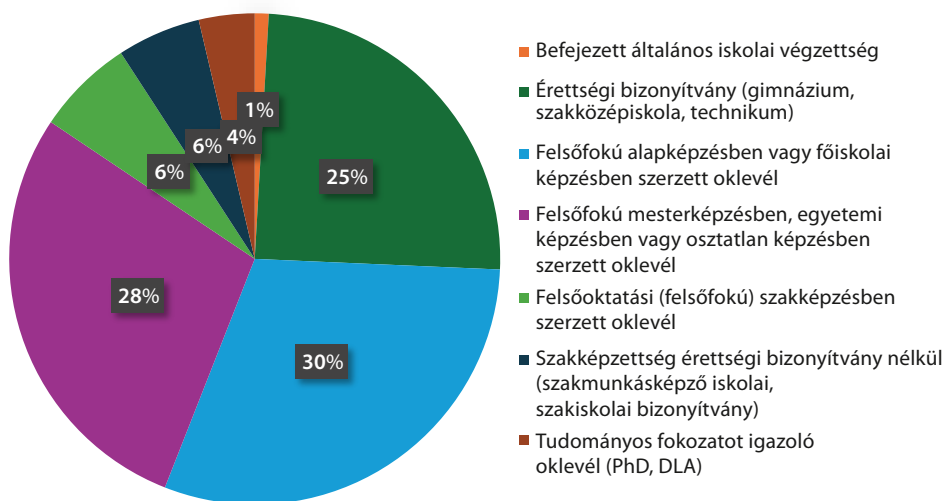
3. ábra: A kérdőívet kitöltők életkori megoszlása

Forrás: a szerző szerkesztése

A nemek megoszlása közel kiegyenlített: a válaszadók 55%-a férfi, 45%-a nő. Ez különösen fontos, mivel a kiberbiztonsági tudatosság vizsgálata során gyakran tapasztalható, hogy a minták nemi alapon torzulnak.

Az iskolai végzettséget tekintve a válaszadók többsége rendelkezik felsőfokú végzettséggel: leggyakoribb kategória a „Főiskolai vagy alapképzésben szerzett diploma”. Ezenkívül jelentős arányban szerepeltek mesterképzéses és középfokú végzettségű válaszadók is (4. ábra), ami lehetővé teszi az iskolai háttér hatásának vizsgálatát az MI-hez való viszonyulás szempontjából.

Mi a legmagasabb iskolai végzettsége?



4. ábra: A kérdőívet kitöltők legmagasabb iskolai végzettsége

Forrás: a szerző szerkesztése

A gazdasági aktivitás tekintetében a legtöbben alkalmazotti jogviszonyban dolgoznak (75 fő, a kitöltők 68%-a), de a mintában szerepeltek diákok, vállalkozók, szülési szabadságon lévők, nyugdíjasok és álláskereső is. Ez a sokféleség támogatja a viselkedései és attitűdbeli különbségek feltárását.

A lakóhely típusa szerint a kitöltők legnagyobb része (40 fő) vármegyeszékhelyen vagy megyei jogú városban él, ami a digitális hozzáférés és a kibertérhez való viszonyulás szempontjából is releváns. Emellett kisebb városok, községek és a főváros is képviseltette magát.

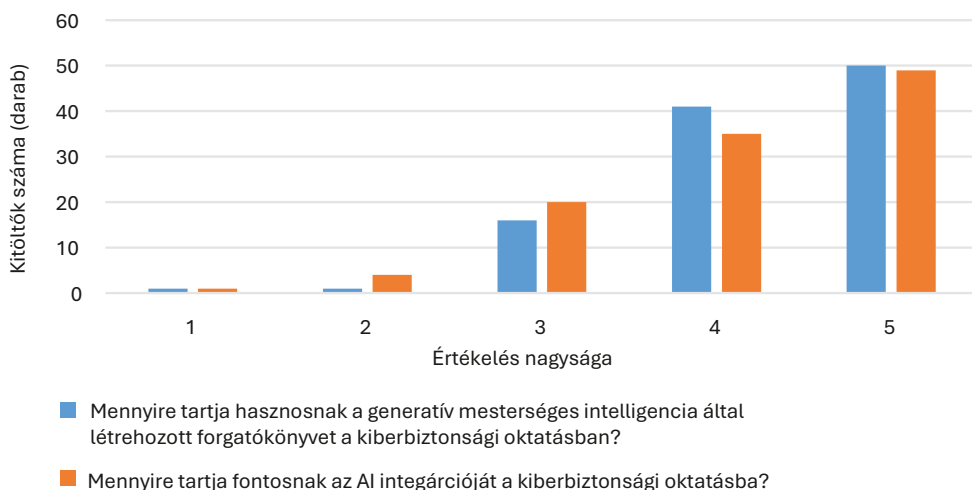
A szakterületi eloszlás igen változatos: 39 különböző megnevezés szerepel. A „Informatika / IT Biztonság” kategóriát 27 fő jelölte meg szakterületi háttérként. Mivel ez a megnevezés nem feltétlenül takar specializált kiberbiztonsági tudást, a tanulmány nem tekinti ezt a csoportot egységes szakértői kategóriának. Ugyanakkor a technológiai affinitás és informatikai háttér releváns szempontként jelenik meg az eredmények elemzése során, így az IT-s háttérrel rendelkező válaszadókat külön bontásban is vizsgálni kell.

A generatív mesterséges intelligencia alapú oktatási tartalmak elfogadottsága

A tanulmány bevezetőjében megfogalmazott első kutatási kérdés azt vizsgálta, hogy a nem szakmabeli felhasználók milyen mértékben fogadják el a generatív mesterséges intelligencia által létrehozott oktatási tartalmakat a kibertudatosítási képzésekben. A kérdéshez kapcsolódó első hipotézis szerint a válaszadók többsége elfogadja a GenAI által generált kiberbiztonsági oktatási szcenáriókat hiteles és hasznos tananyagként.

A kérdőív több szempontból is megvizsgálta az elfogadottság dimenzióját, három különböző kérdéssel, amelyek a TAM komponenseinek is megfeleltethetők. A „Mennyire tartja hasznosnak a generatív mesterséges intelligencia által létrehozott forgatókönyvet a kiberbiztonsági oktatásban?” kérdés az észlelt hasznosságra, míg a „Mennyire tartja fontosnak az AI integrációját a kiberbiztonsági oktatásba?” és az „El tudná képzelni, hogy hasonló interaktív forgatókönyvekkel tanuljon a kiberfenyegetésekről?” kérdések a használati szándékot jelzik.

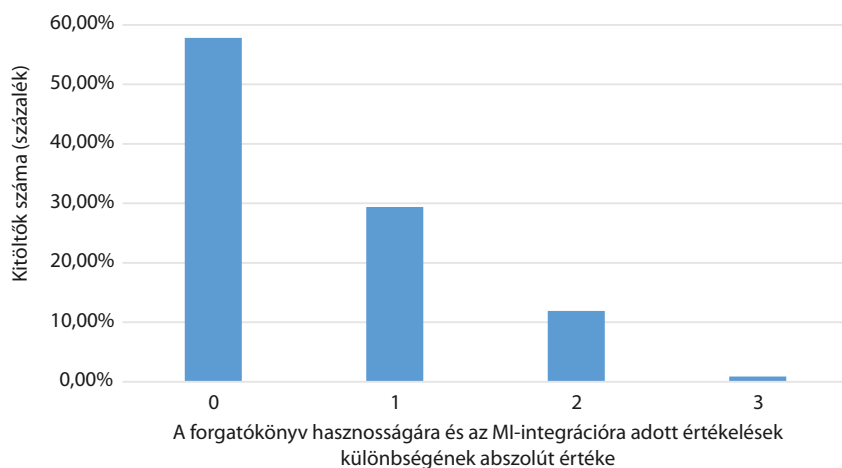
A válaszadás 1-től 5-ig terjedő Likert-skálán történt, a 109 fő által adott válaszok átlaga 4,27, szórása 0,81 volt az esettanulmány hasznosságánál, míg 4,17-os átlag és 0,92-os szórás mutatkozott az AI-integráció fontosságánál. Az eloszlás mindkét esetben enyhén torzult a magasabb értékek irányába, vizuálisan is jól elkülönülő csúcsponttal a 4–5-ös értékeknél (5. ábra).



5. ábra: A kérdőívet kitöltők válasza a fenti kérdésre

Forrás: a szerző szerkesztése

A két kérdés közötti kapcsolat közepesen erős pozitív korrelációt mutatott, azok, akik magasra értékelték a konkrét forgatókönyv hasznosságát, jellemzően az AI-integrációt is fontosnak tartották. A Pearson-féle korrelációs együttható értéke $r = 0,437$ volt ($n = 109$), ami statisztikailag közepesen erős pozitív kapcsolatot jelez a két változó között (6. ábra).

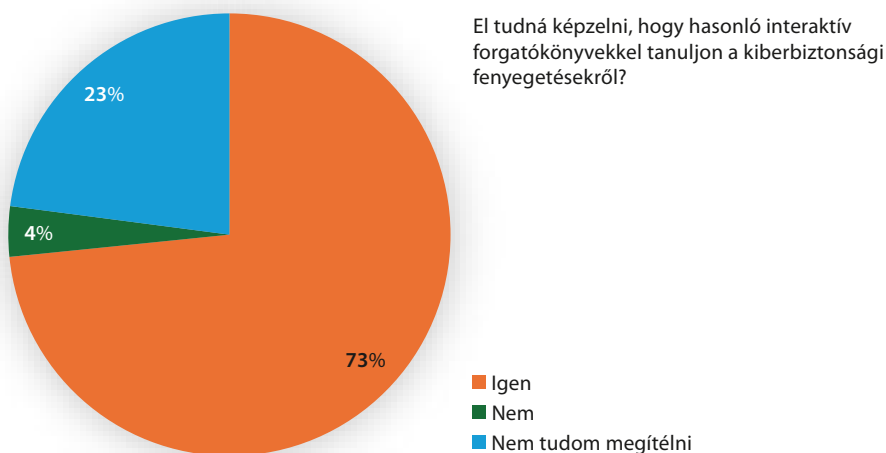


6. ábra: Különbségek eloszlása a két kérdés válasza között

Forrás: a szerző szerkesztése

A diagramon jól látható, hogy az esetek 57,8%-ában a két érték teljesen megegyezett, 29,4%-ban pedig csupán 1 pontos eltérés volt az értékek között. Mindössze 11,9%-nál fordult elő 2 pontos eltérés, és csupán egyetlen válaszadó esetén tért el három vagy több ponttal a két érték. Ez az eredmény jól szemlélteti, hogy a válaszadók nemcsak konzisztens véleményt formáltak, hanem a konkrét tapasztalat és a jövőbeni nyitottság között is szoros kapcsolat figyelhető meg.

A jövőbeni alkalmazhatóság vizsgálatára szolgáló „El tudná képzelni, hogy hasonló interaktív forgatókönyvekkel tanuljon a kiberfenyegetésekről?” kérdésre a 7. diagram szerint érkeztek a válaszok.



7. ábra: Válaszok megoszlása a jövőbeni interaktív forgatókönyveken történő részvételi szándékkal kapcsolatban

Forrás: a szerző szerkesztése

A válaszadók 73,4%-a igennel válaszolt erre a kérdésre, ami egyértelműen arra utal, hogy az MI-alapú tananyagokkal való tanulás nemcsak elfogadható, hanem kifejezetten kívánatos is a felhasználók többsége számára. Ez tovább erősíti a TAM-modell használati szándékra vonatkozó elemét, és azt mutatja, hogy a generatív MI alapú tanulási környezetek a jövőben is beilleszthetők lehetnek a tudatosítási programokba, különösen, ha azok interaktív és szituációalapú formában jelennek meg.

Összességében elmondható, hogy a H1 hipotézist igazoltuk. A válaszadók jelentős többsége hasznosnak találta a GenMI által generált forgatókönyvet, fontosnak tartja az MI integrációját az oktatásba, és nyitott a jövőbeli használatára. Ez az elfogadottság három külön dimenzióban is kimutatható volt: az aktuális tananyag megítélésében (észlelt hasznosság), az elvi támogatottságban (MI-integráció fontossága), valamint a jövőbeni használati szándékban (interaktív szcenáriók alkalmazhatósága). A statisztikai és vizuális eredmények alapján kijelenthető, hogy a generatív mesterséges intelligencia-alapú tananyagok alkalmasak lehetnek a kiberbiztonsági képzések kiegészítésére és fejlesztésére.

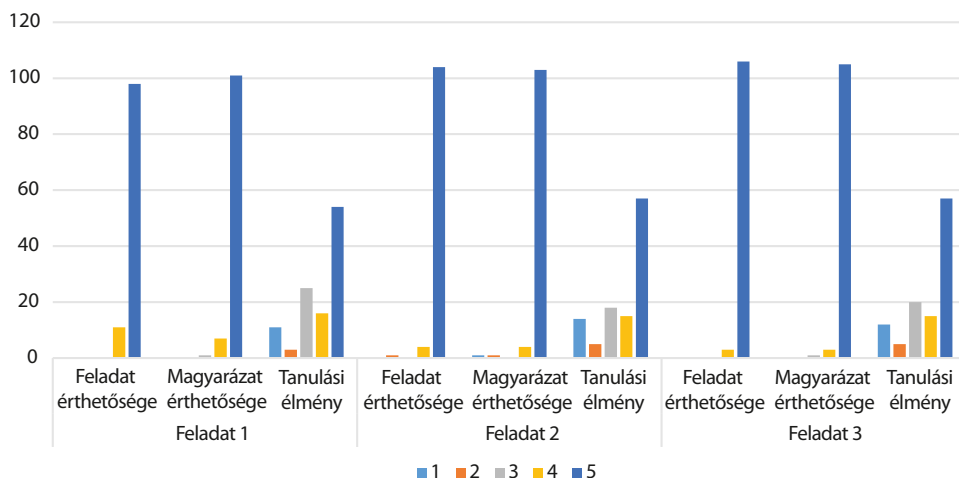
A generatív mesterséges intelligencia tartalmak érthetősége és tanulási értéke

A második kutatási kérdés azt vizsgálta, hogy a válaszadók hogyan értékelik a generatív mesterséges intelligencia által létrehozott kiberbiztonsági tartalmak érthetőségét és tanulási értékét. Az ehhez kapcsolódó második hipotézis szerint a résztvevők az MI által generált tananyagot érthetőnek és tanulásra ösztönzőnek találják.

A kérdőív három egymásra épülő *phishing* feladatot tartalmazott, amelyek fokozatosan növekvő bonyolultságot képviseltek, de nem haladták meg az általános felhasználói tapasztalatokat és technikai képességeket. Minden feladatra előre megadott válaszlehetőségek közül választhattak a kitöltők. A feladatok célja nem a kitöltők által adott válaszok helyességének az ellenőrzése, hiszen azt legjobban nyílt végű kérdésekkel lehetne figyelni, hanem a generatív mesterséges intelligencia által megírt feladat és magyarázat minőségének értékelése.

Minden feladatnál 4 választási lehetőség közül kellett választani a kitöltőnek, a válasz bejelölése és beküldése után azonnali visszacsatolást kapott a kitöltő, ahol a GenMI elmagyarázta, hogy a 4 válasz miért helyes vagy helytelen, valamint egy általános érvényű magyarázatot is adott. Mindhárom visszajelzés után három, azonos szerkezetű értékelő kérdés következett (1–5 skálán): a feladat érthetősége, a magyarázat érthetősége és a tanulás mértéke (8. ábra).

A vizuális eloszlások alapján megállapítható, hogy mind az érthetőség, mind a tanulási élmény többségében magas pontszámokat kapott. Az érthetőség eloszlása enyhén jobbra torzult, a legtöbben 4-es és 5-ös értékeket adtak.



8. ábra: A 3 phishing feladatra adott válaszok eloszlása

Megjegyzés: 1-től 5-ig terjedő skálán, ahol az 1 a legrosszabb, az 5 a legjobb értékelés.

Forrás: a szerző szerkesztése

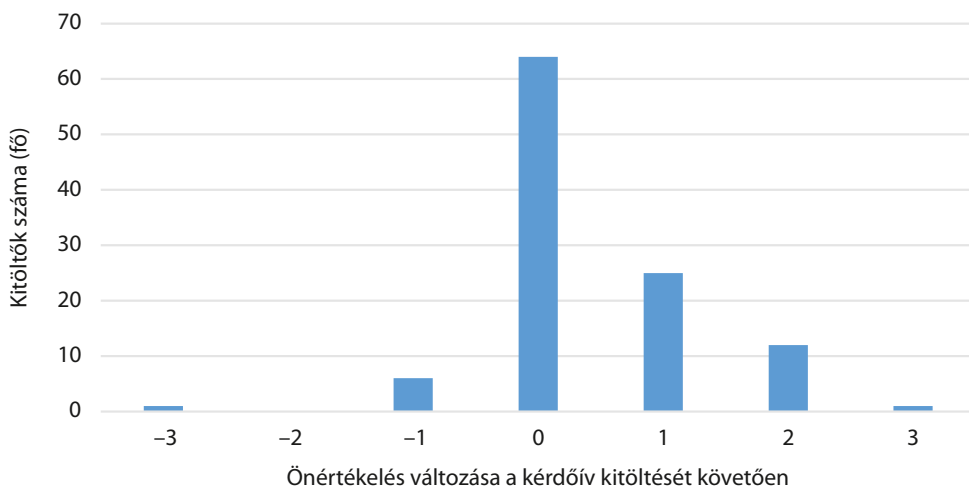
A tanulási élmény eloszlása szélesebb, és bár mediánban hasonló, szórása jóval nagyobb, mint az érthetőségi értékeké. Ez azt jelzi, hogy bár a legtöbb pozitív tanulási élményről számoltak be, az élmény szubjektívebb és egyénibb megítélés alá esik, mint a tartalom érthetősége. Érdeemes megfigyelni, hogy a feladatok fokozatos nehezedeése ellenére az értékelések nem estek vissza jelentősen.

A H2 hipotézis, miszerint a GenMI által generált tartalmak érthetőek és tanulásra ösztönzőek, igazolódott. A válaszadók a *phishing* feladatokat jellemzően jól értelmezhetőnek találták, és egyértelműen jelezték, hogy tanultak is belőlük. Mindez azt mutatja, hogy a GenMI képes a világos, strukturált és tanulást támogató tartalom előállítására nem szakmai célcsoportok számára is.

A tudásszint önértékelésének változása

A harmadik kutatási kérdés arra kereste a választ, hogy a válaszadók saját megítélése szerint mennyiben változott a tudásszintjük a kérdőívben bemutatott *phishing* szcenáriók és feladatok hatására. Az ehhez kapcsolódó harmadik hipotézis szerint a résztvevők az MI-szenárió és feladatok elolvasását követően magasabbra értékelik saját képességüket az adathalász-támadások felismerésében és kivédésében.

A kérdőív két kérdést tartalmazott ezzel kapcsolatban, azonos skálán (1 – egyáltalán nem, 5 – teljes mértékben) mérve az önértékelt jártasságot. Az első kérdés még a GenMI-feladatok és -szcenárió bemutatása előtt szerepelt, míg a második azokat követően. Ez lehetővé tette a közvetlen összehasonlítást, nemcsak statikusan, hanem dinamikus is, az önreflexió szintjén (9. ábra).



9. ábra: Az adathalászat felismerésére irányuló önreflexió változása

Forrás: a szerző szerkesztése

A kapott válaszok különbségének vizsgálatára páros t-próbát alkalmaztunk, amely azt méri, hogy az azonos válaszadók által megadott, szcenárió előtti és utáni értékek közötti eltérés statisztikailag szignifikáns-e. A párosított t-próba eredménye alapján a különbség szignifikáns volt, a $t(108)^{23} = 4,784$, $p < 0,001$, amely azt jelzi, hogy az önértékelt jártasság növekedése statisztikailag szignifikáns, azaz nem magyarázható véletlen ingadozással. Ez megerősíti, hogy a GenMI-feladatok és -szcenárió szignifikáns önértékelési javulást eredményeztek a válaszadók adathalászat elleni jártasságában. Bár az esetek túlnyomó többségében a változás pozitív vagy semleges volt, 7 válaszadónál enyhe csökkenés mutatkozott. Ez feltehetően nem a tudásszint csökkenését, hanem az önreflexió árnyaltabbá válását jelzi. A változatlan válaszadók között 24 fő már eleve maximális, 5-ös értékelést adott magának, ami azt jelzi, hogy az értékelés stabilitása sok esetben nem a hatás hiányát, hanem a meglévő kompetencia érzékelt teljességét tükrözi.

Összességében a H3 hipotézis igazolódott: a válaszadók átlaga szerint a GenMI-feladatokat követően magasabbra értékelték saját jártasságukat az adathalászat felismerésében. Bár a válaszadók több mint fele nem változtatott az értékelésén, a pozitív elmozdulások aránya, valamint az átlagos eltérés mértéke statisztikailag szignifikáns növekedést jelez. Mindez arra utal, hogy a generatív mesterséges intelligencia által előállított tartalom képes volt nemcsak az ismeretek aktiválására, hanem a tudatosság és önreflexió fejlesztésére is – különösen a nem szakmabeli felhasználók körében.

²³ A szabadsági fok ($df = 108$) a páros t-próbára jellemző módon az érvényes válasz párok számából adódik, azaz $df = n - 1$, ahol n azoknak a kitöltőknek a száma, akik mindkét – a szcenárió előtti és utáni – kérdésre válaszoltak.

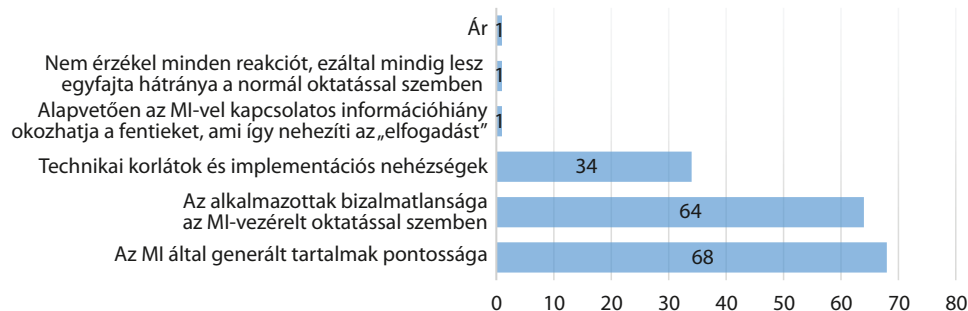
Kihívások a generatív mesterséges intelligencia alapú oktatás bevezetésében

A negyedik kutatási kérdés célja annak feltérképezése volt, hogy a válaszadók szerint mely tényezők jelentik a legnagyobb kihívást az MI-alapú oktatás bevezetése során. A hozzá kapcsolódó H4 hipotézis szerint a leggyakrabban említett kihívás az adat-biztonság és a tartalom megbízhatósága.

A kérdőívben szereplő zárt, többválaszos kérdés pontosan ezekre a tényezőkre kérdezett rá. A válaszadók 62,4%-a tartotta a legnagyobb kihívásnak az MI által generált tartalmak pontosságát, míg 58,7%-a említette az alkalmazottak bizalmatlanságát az MI-vezérelt oktatási rendszerekkel szemben. Ez a két tényező együttesen jól lefedi a H4 hipotézis kulcselemeit, hiszen a válaszok többsége a megbízhatóság és a bizalom hiányát jelöli meg akadályként.

A technikai korlátok és implementációs nehézségek a válaszadók közel egyharmada (31,2%) szerint szintén fontos szempont, de kevésbé meghatározó, mint a tartalmi és emberi tényezők. Az egyéb opciókat (például „információhiány”, „ár”, „nem érzékeli a reakciókat”) csak elenyészően kevesen választották (mindegyik 0,9% alatt), így ezek marginálisnak tekinthetők.

Mit tart a legnagyobb kihívásnak az MI-alapú oktatás bevezetésében?
(Több válasz is választható)



10. ábra: Az MI-alapú oktatás legnagyobb kihívásai

Forrás: a szerző szerkesztése

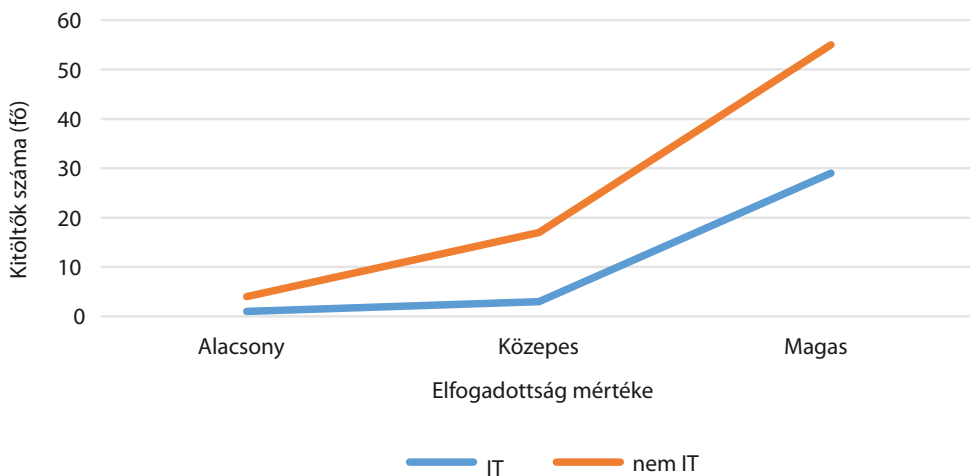
Az eredmények megerősítik azt a feltételezést, hogy a GenMI-alapú oktatás elfogadásának legnagyobb gátja nem technológiai, hanem elsősorban pszichológiai és bizalmi természetű. A válaszok alapján a felhasználók számára az információk pontossága és a rendszerbe vetett bizalom a két legkritikusabb tényező, míg a technikai kihívások másodlagos szerepet játszanak. A H4 hipotézis, miszerint a válaszadók elsősorban a tartalom megbízhatóságát és a bizalom hiányát tartják kihívásnak, igazolódott.

Különbségek generatív mesterséges intelligenciához való hozzáállásban demográfiai bontásban

Az ötödik kutatási kérdés célja annak feltárása volt, hogy a válaszadók MI-hez való hozzáállása eltér-e bizonyos demográfiai szempontok, például szakterület, iskolai végzettség vagy életkor alapján. A kapcsolódó H5 hipotézis szerint az informatikai szakterületen dolgozók szignifikánsan pozitívabb attitűddel viszonyulnak az MI-inTEGRÁCIÓHOZ, mint más területeken dolgozó társaik.

A kérdőív egyik zárt kérdése az MI-alapú oktatás fontosságának megítélésére irányult, amelyben a válaszadók egy ötfokozatú Likert-skálán értékelték, mennyire tartják lényegesnek a mesterséges intelligencia beépítését a kiberbiztonsági képzésekbe. Az elemzés során a válaszadókat szakterület alapján két csoportba soroltuk: informatikai (beleértve az IT-biztonságot) és nem informatikai háttérrel rendelkezők. A két csoport összehasonlításához független mintás t-próbát alkalmaztam. Az informatikai háttérrel rendelkező válaszadók átlaga magasabb volt, mint a más területekről érkezőké, azonban a statisztikai próba nem mutatott ki szignifikáns különbséget a két csoport között. Az eredmény, bár az elméleti elvárásokkal összhangban áll, nem volt elég markáns ahhoz, hogy kizárható legyen a véletlen szerepe.

A kvantitatív vizsgálatot egy keresztábrás khi-négyzet-próbával egészítettem ki, amelyben az MI-elfogadottságot három kategóriára bontottuk: alacsony, közepes és magas szintű elfogadottság. Ezt követően megvizsgáltam, hogy az egyes kategóriákban milyen arányban fordulnak elő informatikai, illetve nem informatikai háttérrel rendelkező válaszadók (11. ábra).



11. ábra: Az MI elfogadottsága az IT- és nem IT-kitöltők között

Forrás: a szerző szerkesztése

A khi-négyzet-próba eredménye megerősítette a t-próba következtetését: nem mutatható ki statisztikailag szignifikáns kapcsolat a szakterület és az MI-elfogadottság

között. A tendencia ugyan megfigyelhető, az informatikai háttérű kitöltők kissé pozitívabban értékelték az MI integrációját, ám a különbség mértéke nem haladta meg a szignifikancia határát.

Összességében a H5 hipotézist, miszerint az informatikai szakterületen dolgozók körében szignifikánsan magasabb az MI-integráció elfogadottsága, nem tudtuk megerősíteni. Bár az eredmények iránya az elvárásokkal összhangban áll, a mintában tapasztalt eltérések nem voltak elég erősek ahhoz, hogy egyértelmű következtetéseket vonjunk le. Az adatok arra utalnak, hogy az MI-hez való pozitív hozzáállás nem kizárólag a szakmai háttér függvénye, hanem valószínűleg összetettebb tényezők, például a technológiai nyitottság, az MI-vel kapcsolatos korábbi tapasztalatok, a munkahelyi kultúra vagy az oktatási módszerekkel szembeni nyitottság is befolyásolják a megítélést.

Összegzés

A kérdőív eredményei megerősítették, hogy a válaszadók alapvetően nyitottak a mesterségesintelligencia-alapú oktatási tartalmak iránt, és pozitívan viszonyulnak a generatív MI alkalmazásához a kiberbiztonsági tudatosságnövelésben. A KK1-re adott válaszok alapján elmondható, hogy a válaszadók túlnyomó többsége hasznosnak találta a GenMI által létrehozott scenáriót, és el tudja képzelni, hogy hasonló módszerekkel tanuljon kiberfenyegetésekről. A H1 hipotézist, miszerint a GenMI-tartalmak elfogadottak és hasznosnak ítélték, megerősítettük.

A KK2 és a H2 hipotézis a generatív mesterséges intelligencia által generált tartalmak érthetőségére és tanulási értékére fókuszált. A három *phishing* feladat kapcsán adott visszajelzésekből világosan kiderült, hogy a résztvevők a feladatokat és magyarázatokat túlnyomórészt jól érthetőnek ítélték, és egyértelműen jelezték, hogy tanultak is belőlük. A tanulási élmény szinte minden résztvevőnél magas szintet ért el, a skálán adott értékek torzítása pedig a pozitív irányba mutatott, megerősítve a H2 hipotézist.

A KK3 az önértékelt tudásszintváltozásra irányult, amelynek vizsgálatához páros mintás t-próbát alkalmaztam. A válaszadók a kérdőív kitöltését követően statisztikailag szignifikánsan magasabb értéket adtak az adathalászat felismerésére vonatkozó jártasságukra, mint a kitöltés előtt. A H3 hipotézis, hogy az MI-scenárió pozitív hatással van az önértékelt tudásszintre, igazolódott.

A KK4 azt vizsgálta, hogy milyen kihívások merülnek fel az MI-alapú oktatás bevezetésével kapcsolatban. A zárt kérdés válaszai alapján a leggyakrabban említett aggályok az MI által generált tartalmak pontossága és a rendszerbe vetett bizalom voltak. Az alkalmazottak bizalmatlansága az MI-vezérelt rendszerek iránt, valamint a technikai és implementációs nehézségek is megjelentek a válaszok között, de kisebb arányban. Ezek az eredmények alátámasztják a H4 hipotézist, amely szerint az MI-alapú oktatás legnagyobb kihívásai a tartalom megbízhatóságával és az emberi tényezőkkel kapcsolatosak.

Végül a KK5 és a H5 hipotézis azt vizsgálta, hogy van-e szignifikáns különbség az MI-integráció elfogadottságában a szakterület, iskolai végzettség vagy életkor szerint. A válaszadók szakterület szerinti bontásában bár megfigyelhető volt egy enyhe

tendencia az informatikai háttérű kitöltők közötti nagyobb elfogadottság felé, sem a független mintás t-próba, sem a khi-négyzet-teszt nem mutatott statisztikailag szignifikáns eltérést. A H5 hipotézist, miszerint az informatikai területen dolgozók körében kimutathatóan magasabb az elfogadottság, nem tudtuk megerősíteni. Az MI-hez való hozzáállás a jelek szerint nem kizárólag a szakmai háttér függvénye, hanem egyéni attitűdök, tapasztalatok és technológiai nyitottság is jelentős szerepet játszanak benne.

A tanulmány célja az volt, hogy feltárja, milyen lehetőségeket rejt a generatív mesterséges intelligencia alkalmazása a kibertudatosítási képzésekben, különös tekintettel a tananyagok személyre szabására, a tanulói élmény fokozására és a viselkedésváltozás támogatására. Az elméleti részben bemutattuk, hogy a GenMI hogyan tudja támogatni az oktatástervezést, például az ADDIE modell egyes fázisaiban, valamint milyen előnyökkel és kockázatokkal jár a tanulói oldalon történő alkalmazása. A vizsgálat figyelmet fordított az etikai, adatvédelmi és elfogadottsági kérdésekre is, amelyek elengedhetetlenek egy ilyen érzékeny területen, mint a kiberbiztonság.

Az empirikus részben egy kérdőíves kutatás segítségével vizsgáltam meg, hogy a magyar válaszadók miként értékelik a GenMI által generált kiberbiztonsági oktatási tartalmakat. A 109 fős minta nem reprezentatív, de kellően változatos volt ahhoz, hogy releváns mintázatokra lehessen következtetni. A kérdéseket a Technology Acceptance Model alapján alakítottam ki, és érintették a tartalom érthetőségét, tanulási értékét, elfogadottságát, jövőbeli alkalmazhatóságát, valamint a válaszadók saját tudásszintjükre vonatkozó önértékelését.

A kutatás főbb eredményei megerősítették a legtöbb hipotézist. A válaszadók többsége hasznosnak találta a GenMI által generált forgatókönyvet, és el tudja képzelni, hogy hasonló módon tanuljon kiberfenyegetésekről (H1). A *phishing* feladatok értékelései alapján a tartalmak érthetők és tanulásra ösztönzőek voltak (H2), és statisztikailag szignifikáns változást figyeltem meg az önértékelt jártasságban is az adathalászat felismerése terén (H3). A válaszadók a legnagyobb kihívásként az MI által generált tartalom megbízhatóságát és az alkalmazottak bizalmát nevezték meg, így a H4 hipotézis is igazolódott. A H5 hipotézis, amely az informatikai területen dolgozók pozitívabb MI-elfogadását feltételezte, nem nyert szignifikáns statisztikai alátámasztást, noha enyhe tendencia megfigyelhető volt.

A tanulmányban megfogalmazott hipotézisek többségét megerősítettük a kvantitatív vizsgálat során. Ezek alapján az alábbi tézisek vonhatók le:

T1: A válaszadók többsége elfogadja a generatív mesterséges intelligencia által létrehozott kiberbiztonsági oktatási tartalmakat hasznos és alkalmazható tananyagként.

T2: A GenMI által létrehozott tartalmakat a résztvevők világosan értelmezhetőnek és tanulásra ösztönzőnek találták.

T3: A GenMI-alapú szcenáriók és feladatok elolvasását követően statisztikailag szignifikáns növekedés figyelhető meg a válaszadók önértékelt jártasságában az adathalászat felismerésében.

T4: A válaszadók szerint a GenMI-alapú oktatás legnagyobb kihívása a tartalom megbízhatósága és az MI-rendszerekbe vetett bizalom hiánya.

T5: A kutatás nem mutatott ki statisztikailag szignifikáns különbséget az MI-integráció elfogadottságában az informatikai és nem informatikai háttérrel rendelkező válaszadók között.

Ezek a tézisek egyértelműen rávilágítanak arra, hogy a generatív mesterséges intelligencia nemcsak technológiai innovációként, hanem oktatási eszközként is releváns, ugyanakkor alkalmazása során elengedhetetlen a tartalom minőségének és a felhasználói bizalomnak a biztosítása. A tanulmányban bemutatott eredmények megalapozzák a jövőbeli kutatások irányait, amelyek vizsgálhatják a GenMI hosszú távú hatásait, a viselkedésváltozást, valamint az MI-alapú tartalom integrálásának lehetőségeit a különböző szervezeti kontextusokban.

A vizsgálat eredményei alapján elmondható, hogy a GenMI alkalmas lehet a kiberbiztonsági képzések támogatására, különösen, ha strukturált, kontextushoz illeszkedő tartalommal és megfelelő visszacsatolással működik. A válaszadók nyitottsága és pozitív visszajelzései alapján a generatív technológia nemcsak technikai lehetőség, hanem oktatásmódszertani innovációként is értelmezhető.

A tanulmány ugyanakkor nem törekszik gyakorlati implementációs javaslatok megfogalmazására, és nem vizsgálta kontrollált környezetben a tanulási hatékonyságot. A jövőbeni kutatások fókuszálhatnak az MI-alapú tartalom hosszú távú hatásaira, a viselkedésváltozás mérésére, valamint annak feltárására, hogy mely szakterületeken és milyen módszerekkel építhető be leghatékonyabban a GenMI a szervezeti oktatásba.

Felhasznált irodalom

- ARIKKAT, Dincy R. et al. (2024): IntellBot: Retrieval Augmented LLM Chatbot for Cyber Threat Knowledge Delivery. Arxiv. Online: <https://arxiv.org/html/2411.05442v1>
- CHANDROTH, Shiju (2025): Microsoft Security Copilot – Paramount is Excited about the most Awaited Generative AI for Cyber Security. Online: <https://paramountassure.com/cloud-security-solutions/microsoft-security-copilot/>
- COLUMBUS, Louis (2023): CrowdStrike Defines a Strong Vision for Generative AI at Fal.Con 2023. Zencoder, 2023. szeptember 21. Online: <https://venturebeat.com/security/crowdstrike-defines-a-strong-vision-for-generative-ai-at-fal-con-2023/>
- CrowdStrike (2025): CrowdStrike Delivers the Next Breakthrough in AI-Powered Agentic Cybersecurity with Charlotte AI Detection Triage. Online: www.crowdstrike.com/en-us/press-releases/crowdstrike-delivers-next-breakthrough-in-ai-powered-agentic-cybersecurity-with-charlotte-ai-detection-triage/
- DAVIS, Fred D. (1989): Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319–340. Online: <https://doi.org/10.2307/249008>
- EU Artificial Intelligence Act. The AI Act Explorer (2025). Online: <https://artificialintelligenceact.eu/ai-act-explorer/>
- HAYDOCK, Walter (2024): Generative AI Risk Cheat Sheet: Training vs. RAG. *Stackaware Blog*, 2024. május 30. Online: <https://blog.stackaware.com/p/ai-training-retrieval-augmented-generation-rag>

- InterSystems (2025): *Retrieval Augmented Generation (RAG): Mi ez és hogyan előzi meg a mesterséges intelligencia hibáit?* Online: www.intersystems.com/hu/resources/retrieval-augmented-generation-rag-mi-ez-es-hogyan-eloz-meg-a-mesterseges-intelligencia-hibait/
- JAKKAL, Vasu (2023): *Introducing Microsoft Security Copilot: Empowering Defenders at the Speed of AI. The Official Microsoft Blog*, 2023. március 28. Online: <https://blogs.microsoft.com/blog/2023/03/28/introducing-microsoft-security-copilot-empowering-defenders-at-the-speed-of-ai/>
- Keepnet (2024): *How Generative AI Is Transforming Security Behavior and Culture Programs*. Online: <https://keepnetlabs.com/blog/how-generative-ai-is-transforming-security-behavior-and-culture-programs>
- REED, Jenn – RAY, Ayan (2024): *CrowdStrike's Charlotte AI – Enhancing Productivity of Cyber Security Analysts with Generative AI Built-on AWS. AWS Blog*, 2024. november 1. Online: <https://aws.amazon.com/blogs/apn/crowdstrike-charlotte-ai-generative-ai-on-aws/>
- TitanHQ (2025): *Hoxhunt and TitanHQ Security Awareness Training*. Online: www.titanhq.com/hoxhunt-vs-titanhq/