

KÖZÖSSÉGIMÉDIA- HASZNÁLAT ÉS BIZTONSÁG

Szerkesztette

Bányász Péter – Kovács Gábor – Veszelszki Ágnes



LUDOVIKA
EGYETEMI KIADÓ

KÖZÖSSÉGIMÉDIA-HASZNÁLAT ÉS BIZTONSÁG:
AZ AKTÍV VÉDEKEZŐ VISELKEDÉSEK
HÁTTÉRTÉNYEZŐI

KÖZÖSSÉGIMÉDIA- HASZNÁLAT ÉS BIZTONSÁG

Az aktív védekező viselkedések
háttértényezői

Szerkesztette

Bányász Péter – Kovács Gábor – Veszelszki Ágnes



LUDOVIKA
EGYETEMI KIADÓ

BUDAPEST, 2026

A tanulmánykötet az NMHH-val kötött, EM-1/1/2022. számú együttműködési megállapodásban részletezett feladatok végrehajtása céljából készült.



Szakmai lektor
Kovács László

Szerzők

Bányász Péter (1., 2., 3., 4. fejezet)
Dub Máté (2., 3., 4., 5. fejezet)
Hankó Viktória (2., 3., 4., 6. fejezet)
Kovács Gábor (1., 5., 6., 7., 8., 9. fejezet)
Laska Pál (2., 3., 4., 7. fejezet)
Mezriczky Marcell (8. fejezet)
Szabó Kincső (9. fejezet)
Veszelszki Ágnes (1., 8., 9. fejezet)

Kiadja a Nemzeti Közszerzői Egyetem
Ludovika Egyetemi Kiadó
A kiadásért felel: Deli Gergely rektor
Székhely: 1083 Budapest, Ludovika tér 2.
Kapcsolat: kiadvanyok@uni-nke.hu · ludovika.hu
Felelős szerkesztő: Karácsony Fanni
Olvasószerkesztő: Fülöp Éva
Tördelőszerkesztő: Fehér Angéla
Korrektor: Bujdosó Hajnalka

ISBN 978-963-653-284-0 (ePDF)
ISBN 978-963-653-285-7 (ePub)
DOI: https://doi.org/10.36250/01289_00

©A szerző, 2026
Szerkesztés © A szerkesztők, 2026

A mű elektronikus változata
a CC BY-NC-ND 4.0 nemzetközi licenc
feltételeinek megfelelően felhasználható.
<https://creativecommons.org/licenses/by-nc-nd/4.0/>



TARTALOM

<i>Előszó</i>	9
Bevezetés	15
1. A felmérés általános adatai	17
A kockázatok tekintetében	17
A kitöltők demográfiai megoszlása	24
2. Védekezés az online visszaélésekkel szemben	33
Technológiai fejlődés és a mesterséges intelligencia	33
A humán tényező szerepe	36
Védekezés az online térben	38
Módszer	41
Eredmények	48
Összegzés	52
3. Az aktív védekezés összefüggései a demográfiai háttérrel	55
Az aktív védekező viselkedés jellemzői és elterjedése	
különböző demográfiai csoportokban	56
Veszélyeztetettebb demográfiai csoportok	
a kiberbiztonság területén	58
Célzott kiberbiztonsági megoldások a veszélyeztetett	
demográfiai csoportok számára	59
Az egyes demográfiai csoportok viselkedési jellemzői	60
Fiatalabb korosztály (18–35 év)	61
Idősebb korosztály (55 év felett)	62
Alacsonyabb gazdasági státuszú felhasználók	63
Gyermekek és tizenévesek	65
Szociálisan elszigetelt vagy elszegényedett közösségek	67

Eredmények	69
Összegzés	74
4. A védelemmotiváció összetevői	77
Módszer: a védelemmotiváció operacionalizációja	80
Eredmények	89
Összegzés	104
5. A demográfiai csoportok közötti különbségek a védelemmotivációban	107
Eredmények	110
Összegzés	116
6. A védelemmotiváció szerepe a védekező viselkedésformák kialakulásában	119
A kiberbiztonsági aspektus	119
Viselkedési formák a kibertérben	120
A védelemmotiváció és a védekező viselkedés kapcsolata	123
Eredmények	126
Összegzés	131
7. A személyiség és védelemmotiváció dimenzióinak összefüggésrendszere	135
A személyiség és a viselkedés összefüggései	135
A Nagy Ötök modell dimenziói és a védelemmotiváció összetevői	136
Módszer	139
Eredmények	144
Összegzés	147
8. A mesterséges intelligencia által generált képek felismerésének képessége	151
Módszer	154
Eredmények	158
Összegzés	167
9. A vizuális MI-felismerő képesség háttértényezői	169
A digitális képmanipuláció korszaka	170
Az MI-felismerő képesség	171

Eredmények	172
Összegzés	179
<i>Felhasznált irodalom</i>	183
<i>A kötet szerzői</i>	193
<i>Függelékek</i>	195
Rövidítések jegyzéke	195
A mesterséges intelligenciával generált képi ingerek előállításához használt promptok	196
A vizsgálatban alkalmazott képek eredeti méretben	200

ELŐSZÓ

A számítógéphez kötődő visszaélések azóta léteznek, amióta a számítógépet feltalálták. Amellett, hogy szórványosan már az internet megjelenése előtt is léteztek számítógép segítségével elkövetett bűncselekmények, ez az internet elterjedésével olyan súlyú probléma lett, hogy az Európa Tanács keretein belül már több mint 20 évvel ezelőtt, 2001-ben – Budapesten – született egy nemzetközi egyezmény, amely a kiberbűncselekményekkel szemben történő egységes nemzetközi fellépést tűzte ki célul. A magyar büntetőjogba is ennek nyomán inkorporálták az ismert számítógépes bűncselekményeket.

Amikor az egyezmény született, az internet még épp csak bontogatta a szárnyait, az internetes bankokat és áruházakat csak egy szűk réteg használta, a közösségi médiának pedig még se híre, se hamva nem volt, a Facebookot is csak három évvel később alapították meg.

Azóta nagyot változott a világ. Ma már alig van ember, akinek ne lenne internetes bankfiókja. A Facebooknak több mint 7 millió felhasználója van csak Magyarországon. Az interneten vásárlunk és tájékozódunk, a pártok legfőbb kommunikációs csatornája a közösségi média lett, és a semmiből jött jelöltek választásokat nyernek egy jól felépített TikTok-kampánnyal. Ahogy az internet átalakult, és ahogy egyre több aktivitásunk költözött a virtuális térbe, úgy változott meg a kiberbűnözés dinamikája is. Eleinte ravas és okos hackerek törtek fel bonyolult számítástechnikai rendszereket, majd a bankkártyákkal való visszaélés jött divatba. Még ma is léteznek persze ezek a „hagyományos” számítógépes bűncselekmények, a furfangos hackerek és a bankkártyacsalások is, de a 2010-es években új kifejezés honosodott meg a kiberbiztonsági szakemberek körében: a *social engineering*. Ez a fogalom olyan új visszaélési formákat takar, amelyek azon az egyszerű megfigyelésen alapulnak, hogy egy komplex számítástechnikai rendszernek a leggyengébb elemei maguk az emberek, akiket viszonylag sablonos, de ugyanakkor a legmélyebb ösztöneikre, félelmeikre, gyengeségeikre ható módszerekkel, *szociális manipulációval* (nem méltányos a mérnököket

belekeverni ebbe) igen hatékonyan lehet irányítani és kihasználni. A könyv 2. fejezetéből (Bányász Péter, Dub Máté, Hankó Viktória, és Laska Pál írása) megtudhatjuk hogy az összes kiberbiztonsági támadás csaknem 30%-a ezen a módszeren alapul, és ezeknek a támadásoknak csaknem 70%-a sikeres is. Magyarországon 2024-ben az MNB adatai szerint csaknem 42 milliárd forint kárt okoztak a banki csalók, az Igazságügyi Minisztérium áldozatsegítő szolgálata szerint hozzájuk 120–160 ilyen jellegű panasz fut be havonta. Szinte mindenkinek van olyan ismerőse, aki már áldozattá vált valamilyen – jellemzően az internetes bankfiókkal kapcsolatos – csalásban. És természetesen a banki csalások csak a jéghegy csúcsát jelentik, a visszaélések ennél jóval szélesebb skálán mozognak.

Ha pedig az ember valóban a leggyengébb láncszem, akkor nagyon fontos kérdés, hogy az emberek hogy tudnak, hogyan akarnak és mennyire képesek védekezni az online térben. A könyv *aktív védekező viselkedésnek* nevezi ezt az attitűdöt, amelyet a szerzőgárda empirikus eszközökkel vizsgált meg. Fontos megjegyezni, hogy a fentebbi, banki csalásokhoz kötődő példákhoz képest a vizsgálat fókusza részben szűkebb, részben azonban tágabb volt. Szűkebb volt, hiszen, ahogy a könyv címe is jelzi, elsősorban a közösségimédia-használat kontextusában vizsgálták az aktív védekező viselkedést, de annyiban tágabb, hogy a könyv nemcsak a kifejezetten csalási, és főleg nemcsak az anyagi hátrányt okozó helyzetekre fókuszál, hanem minden „visszalérésre”, így például a dezinformációra, az anyagi kár nélküli visszaélésekre is, sőt azokra a helyzetekre is, ahol a konkrét „rossz” be sem következik, a felhasználó csak veszélynek van kitéve.

A kötet egy kérdőíves kutatás eredményeit tartalmazza, amelyet a szerzőgárda *Megtévesztések és visszaélések a közösségi média platformjain* címmel végzett a Nemzeti Közszolgálati Egyetem hallgatóinak körében. Magának a kutatásnak a konkrét eredményi is érdekesek, jelen előszó íróját ugyanakkor jobban lenyűgözte az a módszertan, amelynek segítségével a szerzőgárda operacionalizálni igyekezett az aktív védekező magatartást, illetve az ennek mérésére szolgáló kérdéseket. A védekező magatartást a kutatás módszertana két részre, „A” és „B” fődimenziókra bontotta fel. Az „A” fődimenzióhoz „tartoznak mindazon viselkedések, amelyek célja az információ hozzáféréseinek vagy az információ terjedésének hatékony ellenőrzésének, kontrollálásának

biztosítása”, míg a „B” fődimenzió a naprakészséget méri. Mindkét fődimenzión belül két-két aldimenziót találhatunk, az „A”-n belül a személyes adatok védelme és a tudatosan közzétett információk, míg a „B” dimenzión belül egyfelől az ismeretek, másfelől a hardvereszközök naprakészen tartása jelentik a két aldimenziót.

Ahogy ez lenni szokott az empirikus felméréseknél, vannak kevésbé meglepő, mondhatni előrelátható eredmények, de vannak első látásra váratlan adatok is. Talán kevésbé meglepő, hogy a leggyakoribb védekezési módszer az, hogy a felhasználók a privát kommunikációs csatornákat részesítik előnyben a nem privát csatornákkal szemben, míg a legkevésbé az jellemző, hogy elolvassák a közösségimédia-szolgáltatások felhasználási vagy adatvédelmi szabályzatait. Általában is jellemző, hogy a felhasználók az A₂ dimenzióban (tudatos közzététel) óvatosak, míg a B₁ dimenzióban (tájékozódás) nem túl aktívak. Az sem meglepő, hogy a sérülékenység az életkorral együtt nő, mert a tudatosság csökken, vagy hogy a kiberbiztonsági képzés jótékonyan hat a védekező attitűdre.

Kicsit már mélyebb összefüggéseket világít meg a 7. fejezet (Kovács Gábor István és Laska Pál munkája), amely az egyes személyiségdimenziók és a védelemmotiváció (mennyire motivált a felhasználó a védekezésben) összefüggéseit elemzi. De talán itt sem számít nagyon különösnek, hogy az egyébként lelkiismeretes emberek hajlamosabbak aktívabban tenni a saját biztonságuk érdekében.

Ugyanakkor – legalábbis jelen sorok írójának – meglepetést okozott az, hogy hogyan tekintenek a felhasználók az egyes online veszélytípusokra, közelebről mennyire tartják valószínűnek, hogy például mesterséges intelligenciával tévesztik meg őket vagy hekkertámadás áldozatai lesznek (4. fejezet, Bányász Péter, Dub Máté, Hankó Viktória és Laska Pál írása). Nos, e két veszély bekövetkeztét a legtöbben igen valószínűnek gondolják, míg azt, hogy *social engineering* (átverés, szélhámosok, csalás) áldozatai lesznek, igen kevésbé valószínűnek. Ami miatt ez az attitűd meglepő, az az, hogy az empirikus tapasztalatok egyszerűen mást mutatnak: az átverések a leggyakoribb veszélyforrások. Még jobban leegyszerűsítve, a legtöbben attól tartanak a legkevésbé, aminek a legnagyobb a bekövetkezési valószínűsége, mert egyszerűen nem akarják elhinni, hogy ők is átverhetők.

A kötetnek kétségtelenül a legérdekesebb része a 8. fejezet (szerzői Veszelszki Ágnes és Szabó Kincső), amely egy konkrét, manapság forró témának számító területen, a mesterséges intelligencia által generált képeken keresztül mérte meg a felhasználók „becsaphatóságát”. Először is, a felmérés eredménye szintén meglepő bizonyos szempontból, hiszen azt mutatja, hogy a felhasználók sokkal hatékonyabban képesek felmérni a generált képeket, mint azt elsőre gondolnánk. Ezt nyilván az is befolyásolta, hogy a válaszhetőségek úgy voltak megszerkesztve, hogy abban nemcsak az igen – nem tengelyen lehetett válaszolni, hanem kételyt is meg lehetett fogalmazni. Másrészt a kutatás nem állt meg itt, és arra is kíváncsi volt, hogy milyen jelekből következtettek a felhasználók arra, hogy MI-vel generált képekkel állnak szemben. Amellett, hogy a képeken látható személyek életkora is befolyásoló tényező volt (bár az elemszám igen kicsi, kevés képről beszélünk), amennyiben az idősebb embereket ábrázoló képekre jobban gyanakodtak, igen érdekes, hogy a legtöbben a „túl tökéletes” képeket találták gyanúsaknak, ami eszembe juttatja *A tanú* címűfilm örökbecsű *bon mot*-ját: „az a gyanús, ami nem gyanús”. Mivel mindnyájan tudjuk, hogy a technológia fejlődik, és nagy valószínűséggel a felismerési képesség csökkenni fog (mert például ezek a rendszerek megtanulják „elrontani” a képeket), a megoldás ezen a területen nem (csak) a tudatosság és az aktív védekezés fejlesztése lesz, hanem egyszerűen a szintetikus tartalmak kötelező megjelölésének előírása, és ennek vasszigorral történő betartatása.

Utóbbi gondolat már kivezet a témánkból. Az aktív, tudatos védekezés és az ezt fejlesztő oktatás, tudatosítás nagyon fontos dolgok, de ez csak két eleme a nagyobb kiberbiztonsági, sőt inkább általános technológiabiztonsági kirakósnek. A tudatosság növelése és a szigorú jogi szabályok (és persze ezek következetes betartatása) mellett még legalább két elem kell a biztonságos online térhez. A „beépített megfelelés” vagy „beépített védelem” (*compliance és protection by design*) lenne a kirakós harmadik eleme, amely leegyszerűsítve azt jelenti, hogy eleve biztonságos technológiákat kell fejleszteni. A szintetikus tartalmak megjelölésének kötelezettsége nyújtja erre éppen az egyik legjobb példát. Ameddig nem lesz kötelező minden szolgáltató minden szolgáltatásában minden szintetikus outputot letörölhetetlenül

felcímkézni – azaz magába a technológiába beleégetni a megfelelést – addig hiába az oktatás és a jogi szabályok nem sokat segítenek.

Végül a negyedik elem már igazán messzire vezet. A főleg USA-beli és távol-keleti platformszolgáltatások üzleti modellje olykor teljesen ellentétben van az adatvédelmi – és olykor a kiberbiztonsági és emberi jogi – szempontokkal is, hiszen ezek a nagy cégek a minél több személyes adat összegyűjtését, valamint a minél mélyebben bevonódó, minél többet képernyő előtt ülő felhasználót preferálják. Végso soron, amíg ezeknek a nagy cégeknek nem lesz valamilyen, nem erre az üzleti modellre épülő alternatívája, addig ez a probléma soha nem lesz tökéletesen megoldva.

Ez azonban nem változtat azon, hogy ez a könyv egy apró, de fontos hozzájárulás ahhoz, hogy jobban megértsük, hogyan lesz az online tér biztonságosabb.

Budapest, 2025. augusztus. 31.

Zódi Zsolt

BEVEZETÉS

A közösségi média napjaink egyik legmeghatározóbb kommunikációs eszközévé vált, amely alapjaiban formálja az emberek információfogyasztási szokásait, társas kapcsolatait és önállóságát a digitális térben. A közösségi média előnyei mellett jelentős kockázatokat is hordoz, beleértve a dezinformáció terjedését, a digitális biztonsági fenyegetéseket és az egyéni felelősségvállalás hiányát az online térben. Ezen kérdések megértése kiemelt jelentőségű mind az egyének, mind a társadalmi szinten zajló digitális transzformáció szempontjából.

Ez a kötet a Nemzeti Közszolgálati Egyetem és azon tágabb közösség hallgatóinak vizsgálati eredményeit mutatja be, akik a közösségi média használatát és a digitális fenyegetésekkel kapcsolatos védelmi stratégiákat kutatták. A kötet két különálló, de egymással szoros összefüggésben lévő témára összpontosít. Egyrészt a védelemmotivációs elmélet keretrendszerében vizsgálja az egyének képességeit és hajlandóságát arra, hogy felismerjék és reagáljanak az online fenyegetésekre a közösségi média kontextusában, másrészt a generatív mesterséges intelligencia által generált képek felismerésére vonatkozó képességeket tárgyalja. Az itt olvasható írások különös hangsúlyt helyeznek a személyes felelősségvállalás és az önhatékonyság szerepére az online biztonság növelésében, miközben teljesen újszerű eredményeket nyújtanak, amelyek jelentős mértékben hozzájárulnak a közösségi média és a kiberbiztonság metszetében végzett tudományos vizsgálatok fejlődéséhez.

A kiadvány alapját három kulcsfontosságú tanulmány képezi. Nurse,¹ Maples-Keller² és szerzőtársai, valamint Boehmer és szerzőtársai³ módszertani eredményeinek adaptálása egy teljesen új perspektívát nyújt a közösségi média és a kiberbiztonság kapcsolatának feltárásában.

¹ NURSE 2019.

² MAPLES-KELLER et al. 2019.

³ BOEHMER et al. 2015.

A kötetben bemutatott kutatási eredmények a szerzők reményei szerint támogatják az olyan stratégiai kezdeményezéseket, amelyek célja a kiberbiztonsági tudatosság növelése. Meglátásuk szerint a személyes felelősségvállalás hangsúlyozása és a védelmi stratégiákhoz kapcsolódó pozitív megerősítések javíthatják az egyének védelmi magatartását. Eredményeik kiegészítik ezen meglátásokat, különös tekintettel arra, hogyan befolyásolja a közösségi média a fenyegetések észlelését és az arra adott válaszokat.

A kötet friss statisztikai elemzései és empirikus adatai közelebb visznek ahhoz, hogy jobban megértsük, milyen módszerekkel érhető el a hatékony oktatás és prevenció az online térben. Az elemzés során a generatív mesterséges intelligencia által készített képek felismerésének képessége új szempontokat nyújt a kiberbiztonsági oktatás és a dezinformáció elleni küzdelem terén. Ez különösen aktuális, tekintettel a mesterséges intelligenciával kapcsolatos technológiai fejlődésekre és azok társadalmi hatásaira.

A kötet célja, hogy elősegítse a közösségi média kockázataival kapcsolatos kritikus gondolkodást, és hozzájáruljon egy tudatosabb, védekező készségekkel felvértezett felhasználói közösség kialakításához.

1. A FELMÉRÉS ÁLTALÁNOS ADATAI

A közösségi média fokozatosan életünk szerves részévé vált. Jelenléte nemcsak otthonainkban, munkahelyeinken és az iskolákban vált általánossá, hanem mindennapos társadalmi tevékenységeink során is egyre inkább teret hódított: meghatározza a társas kapcsolatainkat, átalakította a kapcsolattartási, informálódási és szórakozási szokásainkat, beszédtemák alapját és platformját is adja egyben. Tekinthetünk a közösségi médiára médiaoptimista vagy -pessimista módon, de az egyértelmű, hogy a digitális térben betöltött meghatározó szerepe olyan társadalmi változásokat indított el, amelyek miatt a hatásai alól nem vonhatjuk ki magunkat.

A közösségi média hasznos eszközt jelent, számos célra használjuk, többek között a szeretteinkkel, barátokkal, munkatársakkal való kapcsolattartásra, üzleti lehetőségeink fejlesztésére, társadalmi és politikai ügyek népszerűsítésére, sőt akár krízishelyzetek kezelésére – legyen szó katasztrófa idején történő gyors információközlésről vagy segítségnyújtásról. Azonban mint minden új technológiának, a közösségi médiának is két oldala van: a pozitív hatások mellett számos kihívást és kockázatot rejt, amelyek jelentős mértékben érintik a honvédelmi, rendvédelmi és nemzetbiztonsági szervezetek működését. E szervek működésére is kettős hatása van: egyrészt veszélyeztet(het)i a(z információ)biztonságot, másrészt új lehetőségeket kínál a törvényben meghatározott feladatok hatékonyabb ellátására.

A KOCKÁZATOK TEKINTETÉBEN

A közösségi média kockázatait fokozza, hogy széles körben hozzáférhető különböző eszközökön – többek között PC-ken, tableteken, okostelefonokon, okoszemüvegeken, okostelevíziókon és okosórákon keresztül –, amelyek különböző szintű és jellegű sebezhetőségeket mutatnak. Ez a többszörös hozzáférési pont nehezíti a biztonsági szabályozások kialakítását, hiszen

az eltérő eszközök védelme különböző technikai megoldásokat igényel. Továbbá a hordozható eszközök lehetőséget biztosítanak a munkahelyi tiltások kijátszására, például a munkavállalók saját eszközeik használatával akár tiltott közösségi oldalakhoz is hozzáférhetnek munkaidő alatt.

A közösségi média használata a munka és magánélet közötti határok is elmossa, különösen a szellemi munkát végzők esetében. Nemcsak a munkahelyen élünk társasági életet a közösségi oldalak révén, hanem sok esetben otthon is dolgozunk, ami gyakran háttérbe szorítja a családi programokat és szabadidős tevékenységeket. A technológiai innováció, valamint a közösségi platformok folyamatos átalakulása további kihívások elé állítja a szabályozókat, hiszen úgy kell megalkotniuk az adat- és információbiztonságot garantáló előírásokat, hogy a szabályozási és szabályozandó környezet maga is folyamatos változásban van.

Mindez azt is jelenti, hogy nemcsak a védelmi intézkedéseknek kell folyamatosan naprakésznék lenniük, hanem azoknak a módszereknek is, amelyekkel a közösségi média lehetőségeit a szervezeti feladatellátás szolgálatába állítják. A gyorsan fejlődő technológia és a folyamatosan változó platformok új lehetőségeket teremtenek, ugyanakkor régebbi eljárásokat elavulttá tehetnek. Ezért különösen fontos, hogy az oktatás is lépést tartson e változásokkal, ez a terület véleményünk szerint fejlesztésre szorul annak érdekében, hogy a felhasználók és a szakemberek egyaránt megfelelő tudatossággal és képességekkel rendelkezzenek a közösségi média világában rejlő lehetőségek és veszélyek kezelésére.

A közösségi média kiberbiztonsági kockázatainak részletes ismertetése meghaladja e kötet terjedelmi korlátait, hiszen az elmúlt években számos hazai és nemzetközi tanulmány dolgozta fel alaposan e kérdéskört. A közösségi média nem megfelelő használata számos kiberbiztonsági kockázatot hordoz magában, amelyek mind az egyének, mind a szervezetek számára komoly fenyegetést jelentenek. Az egyik legjelentősebb veszély az adathalászat (*phishing*), amely során a támadók hamis üzenetek vagy weboldalak révén próbálják megszerezni a felhasználók érzékeny adatait, például jelszavakat vagy banki információkat. A közösségimédia-platformok gyakran szolgálnak ilyen támadások kiindulópontjaként, mivel a felhasználók hajlamosak megbízni az ismerősöktől érkező üzenetekben, így könnyebben válnak áldozattá.

Továbbá a nem megfelelő adatvédelmi beállítások lehetővé teszik a támadók számára, hogy személyes információkat gyűjtsenek a felhasználókról, amelyeket később célzott támadásokhoz, például *social engineering* technikához használhatnak fel. Ezen túlmenően, a közösségimédia-platformokon terjedő rosszindulatú szoftverek (*malware*) is jelentős kockázatot jelentenek, amelyek révén a támadók hozzáférhetnek a felhasználók eszközeihez, adatlopást vagy rendszerek feletti irányítást eredményezve. Ezen túl a közösségi média nem megfelelő használata a szervezetek számára is kockázatokat rejt, mivel az alkalmazottak által megosztott információk révén a támadók értékes adatokat szerezhetnek a vállalati infrastruktúráról, működésről vagy biztonsági intézkedésekről. Ez különösen igaz a kritikus infrastruktúrákat üzemeltető szervezetek esetében, ahol a közösségi médián keresztül kiszivárgott információk súlyos következményekkel járhatnak.

A közösségi média jelentős kockázatokat hordoz a dezinformáció terjesztésében, ami mélyreható hatással lehet mind a társadalom egészére, mind az egyének döntéshozatalára. Az egyik legfőbb probléma abból fakad, hogy a közösségi platformok algoritmusai előnyben részesítik a figyelemfelkeltő, akár polarizáló tartalmakat, hiszen a céljuk a felhasználói elkötelezettség állandó növelése. Ez a gyakorlat közvetett módon kedvez a félrevezető információk terjedésének, mivel ezek az üzenetek könnyebben elérik és befolyásolják a felhasználók széles körét, különösen a társadalompolitikai és egészségügyi válságok idején.⁴

A dezinformáció veszélye különösen fokozott olyan időkben, amikor a pontos és megbízható információk nélkülözhetetlenek, például választási időszakokban vagy közegészségügyi krízisek során. Ekkor az álhírek és torzított narratívák gyorsan deformálhatják a valóság észlelését, sőt, közvetlenül befolyásolhatják a közösségi attitűdöket és viselkedési mintákat. A közösségi média eszközként szolgál arra, hogy az érzelmi töltetű, gyakran félelmet vagy haragot kiváltó információk gyors megosztást nyerjenek, ami nagyban megkönnyíti a dezinformáció térnyerését és annak valóságként való elfogadását.⁵

⁴ CINELLI et al. 2020.

⁵ VOSOUGHI–ROY–ARAL 2018.

Ezen túlmenően a botok és egyéb automatizált fiókok által végzett manipulációk szintén kulcsfontosságú tényezők a közösségi média ökoszisztémájában. Az ilyen automatizált szereplők képesek a dezinformációs tartalmak sokszorosítására és felerősítésére, megtevesztve a felhasználókat azáltal, hogy mesterséges közvéleményt generálnak. Ez nemcsak a társadalmi megosztottságot fokozza, hanem jelentős biztonsági kockázatot is jelent, mivel alapjaiban rengetheti meg a demokratikus intézmények iránti bizalmat, valamint fokozhatja az államok közötti politikai és társadalmi feszültségeket.⁶

A közösségi média kiberbiztonsági kockázatait az elmúlt évek technológiai fejlődése, ami elsősorban a generatív mesterséges intelligencia (GenMI) elterjedésében azonosítható, tovább növelte kitettségünket a kibertérből érkező fenyegetésekkel szemben. Az egyik legjelentősebb veszélyforrás az, hogy a GenMI-technológiák segítségével a támadók egyre kifinomultabb és hitelesebb adathalászati (*phishing*) kampányokat tudnak létrehozni. A generatív modellek képesek természetes nyelven megszólaló, személyre szabott üzeneteket vagy hamis identitásokat létrehozni, amelyekkel a felhasználók bizalmát könnyebben megszerzik, így nagyobb eséllyel jutnak hozzá érzékeny adatokhoz.⁷

A GenMI-technológiákat a dezinformáció terjesztésére is felhasználhatják. A hamis hírek és propagandaanyagok automatikus generálása és terjesztése révén a támadók manipulálhatják a közvéleményt, ami különösen veszélyes lehet választási kampányok vagy társadalmi konfliktusok idején. Egy másik kockázatot jelent az úgynevezett *deepfake* technológia, amely lehetővé teszi hitelesnek tűnő képek, hangfelvételek vagy videók létrehozását, ezáltal visszaélésekhez vezethet, például hamis videók használatával történő zsarolás vagy hitelrontás céljából.⁸

A generatív mesterséges intelligencia kockázatai közé sorolhatjuk az adatbiztonsági aggályokat is, mivel ezek a modellek gyakran nagy méretű adathalmazokon tanulnak, amelyeket nem mindig védenek megfelelően. Az ilyen modellek rosszindulatú hozzáférés esetén visszafejthetők, így

⁶ FERRARA 2017.

⁷ HANCOCK–NAAMAN–LEVY 2020.

⁸ MIRSKY–LEE 2021.

érzékeny adatok kerülhetnek nyilvánosságra vagy illetéktelen kezekbe, de ezek rossz, hamis adatokkal való fertőzése is bekövetkezhet.⁹

A generatív mesterséges intelligencia modellek fejlesztése során az egyik legjelentősebb kihívás az alkalmazott adatok minősége és integritása. Amennyiben az adatok megbízhatósága nem garantálható, fennáll a veszélye annak, hogy a modellek tanítása során mesterségesen generált tartalmakat használnak fel, ami torzíthatja a tanulási folyamatot és csökkentheti a modellek teljesítményét. A *blockchain* technológia alkalmazása lehetőséget nyújt az adatok integritásának biztosítására, mivel decentralizált és átlátható módon tárolja az információkat, ezáltal megakadályozva azok manipulálását. A *blockchain* és a mesterséges intelligencia integrációja új perspektívákat nyit az adatok biztonságos kezelésében és a modellek megbízhatóságának növelésében.¹⁰

A jelen kötet alapjául szolgáló kutatásunk egyik alapvető kérdésköre az adat- és információbiztonsági tudatosság, amelyben kiemelt szerepet tulajdonítunk a közösségi média használatának, illetve a mesterséges intelligencia által generált képek felismerésével kapcsolatos képesség vizsgálatának.

Úgy véljük, a digitális immunitás kialakításában az oktatásnak különösen fontos szerepe van. A megfelelő kiberbiztonsági készségek megtanulása ma már nemcsak a szakemberek, hanem mindenki számára fontos, hiszen a digitális világ fenyegetései folyamatosan változnak és egyre kifinomultabbak. Az oktatásnak kulcsszerepe van abban, hogy a hétköznapi emberek is megismerjék az alapvető biztonsági szabályokat és eszközöket, amikkel megvédhetik magukat online. Az ilyen jellegű tudás megszerzése nemcsak a technikai alapokat foglalja magában, hanem fejleszti az emberek tudatosságát is, hogy miként azonosítsanak gyanús tevékenységeket vagy üzeneteket. A kiberbiztonsági képzések során például megtanulható, hogyan védjük meg személyes adatainkat, hogyan válasszunk biztonságos jelszavakat, és hogyan ismerjük fel a csalásokat vagy adathalász próbálkozásokat. Az ilyen képzések révén mindenki felkészültebben és magabiztosabban tud mozogni az online térben, így hatékonyabban védheti meg magát és családját a digitális fenyegetésekkel szemben. A kiberbiztonság területén

⁹ CARLINI et al. 2020.

¹⁰ ESZTERI 2020.

hatékony oktatási módszerek közé tartoznak a gyakorlati laborok, szimulációs gyakorlatok, valamint a valósághú környezetben zajló *capture the flag* (továbbiakban: CTF) versenyek, amelyek hozzájárulnak a jövőbeli szakemberek készségeinek fejlesztéséhez. Švábenský és szerzőtársai 2021-es tanulmánya szerint a CTF-versenyek hatékonyan fejlesztik a kiberbiztonsági ismereteket és készségeket, különösen a kriptográfia és a hálózati biztonság területén, ugyanakkor rámutatnak arra is, hogy az emberi tényezők, mint például a *social engineering* és a kiberbiztonsági tudatosság, gyakran háttérbe szorulnak.¹¹ A szerzők hangsúlyozzák, hogy a jövőbeli CTF-versenyeknek integrálniuk kellene a nem technikai aspektusokat is, hogy jobban felkészítsék a résztvevőket a modern kiberfenyegetésekre.

Megítélésünk szerint a Nemzeti Közzolgálati Egyetem (továbbiakban NKE) hallgatói ezen a téren különösen releváns célcsoportot alkotnak, hiszen az intézmény elvégzését követően várhatóan az állami szféra különböző szakterületein fognak elhelyezkedni. Ha hallgatóink nem kapnak megfelelő adat- és információbiztonsági képzést, amely lehetővé tenné számukra a potenciális fenyegetések felismerését és elhárítását, az őket alkalmazó szervezetek számára jelentős kockázatot jelenthetnek. Az ilyen szervezetek jogszabályban rögzített feladataik ellátása során érzékeny adatokat kezelnek, amelyek nemcsak idegen nemzetek nemzetbiztonsági szervei, hanem kiberbűnözők, terroristák és hacktivisták csoportok számára is értékesek lehetnek.

A hírszerzés történetében nem újdonság, hogy idegen államok nemzetbiztonsági szolgálatai a pályájuk elején álló, várhatóan fontos beosztásba kerülő személyeket tartják beszerzésre érdemesnek. Egy egyetemi éveit töltő fiatal gyakran kevésbé elővigyázatos, és nem feltétlenül tulajdonít kellő jelentőséget az adat- és információbiztonsági szempontoknak, így róluk több kompromittáló információ gyűjthető össze. E típusú beszerzés klasszikus példája a cambridge-i ötök esete, amely az 1930-as években kezdődött, amikor Alexander Mihajlovics Orlov ezredes sikeresen toborozta Kim Philbyt és társait a szovjet hírszerzés számára.¹² Ebben az időszakban

¹¹ ŠVÁBENSKÝ et al. 2021.

¹² HEFFERNAN 2015.

olyan tényezők is közrejátszottak a beszerzés sikerében, mint az érintettek szexuális orientációja, ami abban az időben Nagy-Britanniában komoly társadalmi megvetés tárgyát képezte, így ezek az információk zsarolási alapot is jelenthettek. A Cambridge-ben végzett hallgatók jellemzően magas pozíciókat töltenek be a politikai, katonai és gazdasági szférában, ahogy az Philby esetében is történt, aki az MI-6 kötelékében haladt előre, mígnem a Központi Hírszerző Ügynökség és az MI-6 közötti összekötő tiszti pozícióba került, ezzel a szovjet hírszerzés számára mind a brit, mind az amerikai titkokhoz hozzáférést biztosított.

Az oktatásban az utóbbi évek során számos változás figyelhető meg, és oktatóként, függetlenül az oktatott tantárgytól, közös kihívásokkal szembesülünk. Az ismeretanyag átadása során olyan pedagógiai módszereket szükséges alkalmaznunk, amelyek képesek hosszabb ideig fenntartani a hallgatók figyelmét, hiszen az okoseszközök – különösen a mobiltelefonok – használata gyakran elvonja a figyelmüket. Ezt a helyzetet tovább súlyosbítja, ha a kurzus az esti órákban zajlik, amikor a hallgatók egy teljes napos tanulási és előadás-látogatási sorozat után már fáradtak, ami jelentősen csökkentheti a koncentrációjukat és a tananyag iránti érdeklődésüket. Még azok számára is, akik érdeklődést mutatnak a téma iránt, a mentális leterheltség akadályozhatja az órai tartalmak hatékony befogadását és feldolgozását. Emellett gyakran előfordul, hogy egyes hallgatók kizárólag azért vesznek részt az adott kurzuson, mert a diploma megszerzéséhez szükséges krediteket szeretnék teljesíteni, és ezáltal a tárgy iránti valódi érdeklődés hiánya tovább nehezíti az oktatási célok megvalósítását.

A fenti tanulságok, véleményünk szerint, a Nemzeti Közzolgálati Egyetem hallgatóira is érvényesek. E felismerés alapján végeztünk egy kérdőíves felmérést az egyetem hallgatói körében, amelynek célja az adat- és információbiztonsági tudatosság mértékének feltárása volt. Az eredmények alapján levonható következtetéseket az adat- és információbiztonsági oktatás fejlesztésében kívánjuk felhasználni, hogy hallgatóink képesek legyenek megfelelően kezelni a jövőbeli munkájuk során felmerülő biztonsági kihívásokat.

A kérdőíves felmérésünket *Megtévesztések és visszaélések a közösségi média platformjain* címen szerkesztettük meg, és a Nemzeti Közszerológati Egyetem hallgatói körében vártuk a kitöltéseket az egyetem itt ismertetett, speciális feladatrendszere okán. Meglepetésünkre végül a kitöltők között nem kizárólag NKE-s hallgatókat találhattunk.

A kérdőívünk öt fő egységből állt, és kitöltése nagyjából 30 percig tartott. Az időkeret azért fontos, mert a generatív mesterséges intelligencia felismerésével kapcsolatos képesség vizsgálata során a kitöltőket nem helyeztük időnyomás alá, annyi ideig nézhették, vizsgálhatták a képeket, ameddig szerették volna.

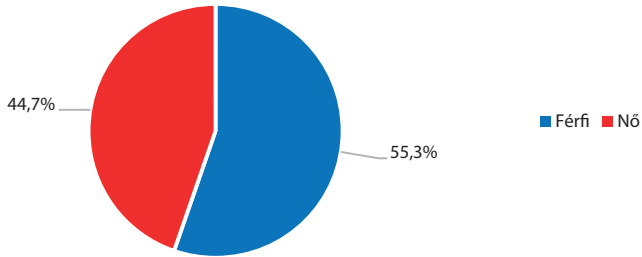
A kérdőív kitöltése természetesen anonim volt, a kitöltőktől semmilyen személyes adatot nem gyűjtöttünk a GDPR-nek történő megfelelés érdekében. A kérdésekre adott válaszokat az SPSS statisztikai program segítségével értékeltük ki, amely személyazonosításra nem alkalmas – azaz anonim – módon tárolja a válaszokat egy adatbázisban, azok összekapcsolása a válaszadó személyével egyéb módon sem lehetséges.

A kérdőív szerkesztése során alkalmazott módszertant a kötet vonatkozó fejezetei részletesen ismertetik.

A KITÖLTŐK DEMOGRÁFIAI MEGOSZLÁSA

A kérdőív kitöltésére négy hónap állt a hallgatók rendelkezésére 2024. június és szeptember között. Ezen időszak alatt összesen 216 kitöltés érkezett, de egy kitöltő által adott válaszokat eltávolítottuk, ugyanis ő életkornak 120 évet írt be, így feltételeztük, hogy a további kérdések megválaszolása is hasonló komolysággal történt.

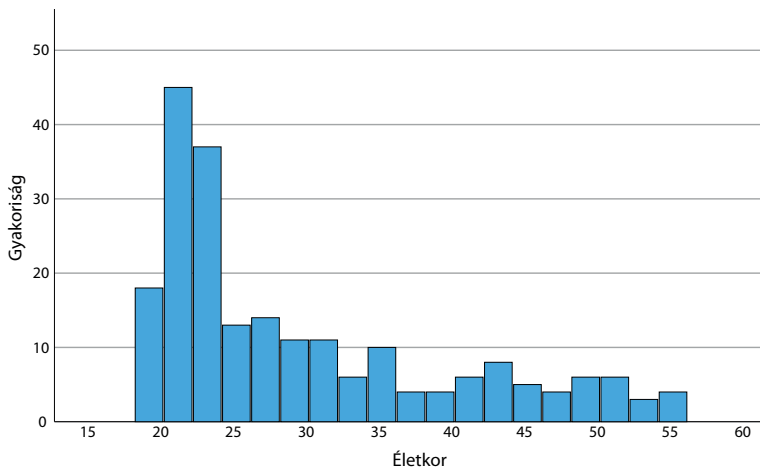
A kitöltők nemek szerinti megoszlása az 1.1. számú ábrán látható, a minta nemi megoszlása közel egyenletes, azonban a férfiak kissé nagyobb arányban vannak jelen (119 férfi, 96 nő).



1.1. ábra: A kitöltők megoszlása nemek szerint (N = 215)

Forrás: saját szerkesztés

Az 1.2. számú ábrán tüntettük fel a kitöltők korosztály szerinti megoszlását.

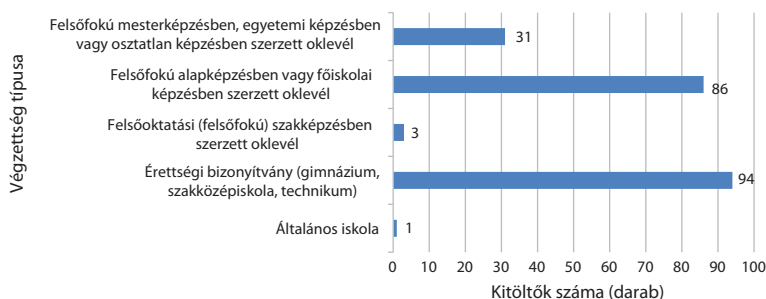


1.2. ábra: A kitöltők megoszlása nemek szerint (N = 215)

Forrás: saját szerkesztés

Az életkor megoszlását ábrázoló hisztogram alapján megfigyelhető, hogy a mintát alkotó egyetemi hallgatók életkori eloszlása a fiatalabb korcsoportok irányába tolódik. A válaszadók többsége a 18 és 25 év közötti korcsoportokhoz sorolható, amely az egyetemi hallgatók körében jellemző életkori intervallumnak tekinthető. A 20–25 éves korcsoport kiemelkedő gyakorisággal jelenik meg a mintában, ami azt jelzi, hogy ebben a korban található a hallgatók többsége, ezáltal jól reprezentálja az egyetemi korosztályt. Az idősebb korcsoportokban, különösen 30 év felett, a válaszadók száma jelentősen csökken, és az 50 év felettiek aránya elenyésző, ami egyértelmű lefelé mutató trendet jelez az életkor növekedésével. Ez az eloszlás természetesnek tekinthető az egyetemi hallgatók mintájában, ahol a résztvevők többsége hagyományosan a fiatal felnőttek közül kerül ki. Megállapítható, hogy ez az eloszlás nem torzított, hanem a hallgatói populáció sajátosságait mutatja, és ennek megfelelően a kapott eredmények relevánsak és jól általánosíthatók az egyetemi hallgatók körében.

Az iskolai végzettség megoszlásának elemzése alapján megállapítható (1.3. ábra), hogy a minta válaszadói között az érettségi bizonyítvánnyal (gimnáziumi, szakközépiskolai vagy technikumi végzettséggel) rendelkezők vannak legnagyobb arányban, akik a teljes minta 43,7%-át teszik ki, 94 kitöltővel. Ez a csoport képezi a legnagyobb részét a hallgatói mintának, ami valószínűsíthetően annak köszönhető, hogy ők az egyetemi képzések alapszakos évfolyamainak hallgatói.

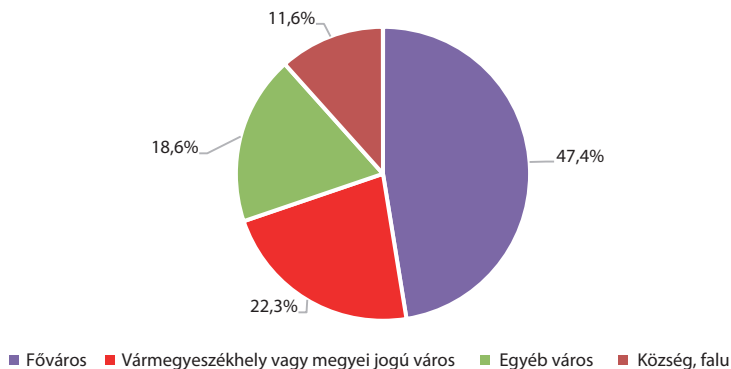


1.3. ábra: A kitöltők iskolai végzettségének megoszlása ($N = 215$)

Forrás: saját szerkesztés

A válaszadók 40,0%-a (86 fő) rendelkezik felsőfokú alapképzésben vagy főiskolai képzésben szerzett oklevéllel, ami azt jelzi, hogy a mintában jelentős számban vannak olyanok, akik már egyetemi vagy főiskolai végzettséggel rendelkeznek, esetleg már befejezték az alapképzési szintet. Ez a csoport a második legnagyobb arányú a válaszadók között, és az egyetemi hallgatók különböző képzési szintjeit tükrözheti. A mesterképzésben, egyetemi képzésben vagy osztatlan képzésben szerzett oklevéllel rendelkezők aránya 14,4%, vagyis 31 fő, ami azt jelzi, hogy a mintában kisebb, de számottevő részarányt képviselnek azok, akik már a felsőoktatás magasabb szintjén járnak, vagy már elvégezték azt. A legkisebb arányban vannak jelen azok, akik csak általános iskolai végzettséggel (1 fő) vagy felsőoktatási szakképzésben szerzett oklevéllel (3 fő) rendelkeznek. A korábbiakban utaltunk már arra, hogy a kérdőívet tervezetten az NKE hallgatói körében töltöttük ki, viszont meglepve tapasztaltuk, hogy nem kizárólag az egyetem hallgató körében töltötték ki, ez magyarázza az 1 fő általános iskolai és 3 fő felsőoktatási szakképzéssel szerzett végzettséggel rendelkező kitöltő válaszait. A későbbiekben ennek a 4 kitöltőnek a válaszait nem elemeztük az iskolai végzettséggel összefüggésben.

Az 1.4. számú ábrán a kitöltők lakóhely szerinti megoszlását láthatjuk.



1.4. ábra: A kitöltők lakóhely szerinti megoszlása ($N = 215$)

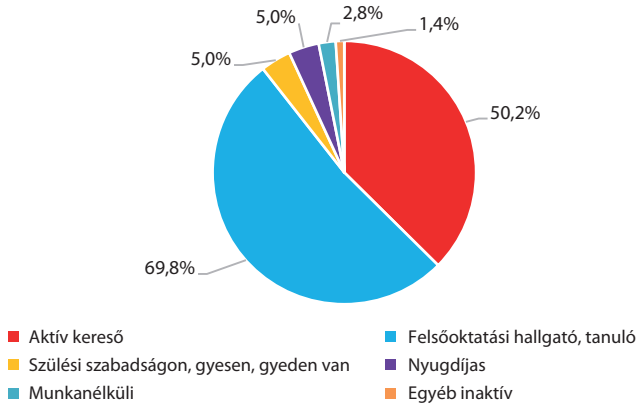
Forrás: saját szerkesztés

Az adatok azt mutatják, hogy a kitöltők közel fele, pontosabban 47,4%-a a fővárosban él, ami kiemelkedő arányt jelent a mintán belül. A második legnagyobb csoportot a vármegyeszékhelyeken vagy megyei jogú városokban élők alkotják, akik a válaszadók 22,3%-át teszik ki. A válaszadók 18,6%-a egyéb városokból származik, ami azt jelzi, hogy a kisebb városokban élők jelenléte is számottevő, de kevésbé meghatározó, míg a mintában a legkisebb csoportot a községekben és falvakban élők alkotják, akik a válaszadók 11,6%-át képviselik.

A foglalkoztatási csoportokra vonatkozó kérdések során a kitöltők több kategóriát jelölhettek (lásd 1.5. ábra). A válaszok érvényességét keresztábra-analízissel ellenőriztük, és nem állapíthattunk meg anomáliákat az eredmények során.

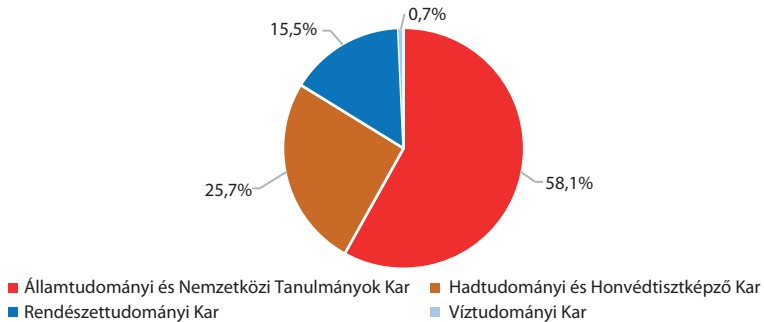
A minta túlnyomó részét, 69,8%-át „Felsőoktatási hallgató, tanuló” kategóriába tartozó személyek alkotják, ami megfelel a vizsgálat célcsoportjának. Az „Aktív kereső” kategóriába tartozók aránya szintén magas, 50,2%, ami arra utal, hogy sok hallgató már tanulmányai alatt is dolgozik vagy levelező képzési formában folytatja tanulmányait. Ez a magas érték azt sugallja, hogy a hallgatók között nagy az igény a munkavállalásra, akár részmunkaidős, akár rugalmas foglalkoztatási formák keretében, amelyek lehetővé teszik a tanulmányokkal való összehangolást. A válaszadók kisebb része olyan kategóriákban található, mint a „Születési szabadságon, gyesen, gyeden lévők” (2,8%), „Egyéb inaktívak” (1,4%), illetve 1-1 kitöltővel a „Nyugdíjasok” (0,5%) és a „Munkanélküliek” (0,5%).

Mint többször említettük, alapvetően a Nemzeti Közzolgálati Egyetem hallgatói körében kívántuk elvégezni a felmérésünket, de a 215 kitöltőből csupán 148-an nyilatkoztak úgy, hogy jogviszonnyal rendelkeznek az egyetemen. E kötet írásakor az NKE-n öt kar lát el oktatási tevékenységet, azonban a kérdőív megszerkesztésekor és a kitöltési időszakban csupán négy kar működött. A Nemeskürty István nevét viselő tanárképző kar 2024. augusztus 1-jén jött létre a Nemzeti Közzolgálati Egyetemen, így a válaszok csupán a korábbi négy kar között oszlanak meg (lásd 1.6. ábra).



1.5. ábra: A kitöltők foglalkoztatási csoport szerinti megoszlása ($N = 215$)

Forrás: saját szerkesztés



1.6. ábra: A Nemzeti Közszolgálati Egyetem hallgatóinak karonkénti megoszlása ($N = 148$)

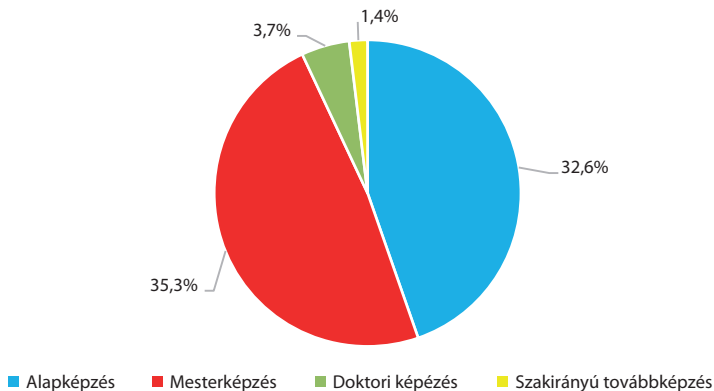
Forrás: saját szerkesztés

A válaszadók kari megoszlásának elemzése alapján a résztvevők túlnyomó része az Államtudományi és Nemzetközi Tanulmányok Karról érkezett, akik a teljes minta 40%-át alkotják, és az NKE hallgatói közül 58,1%-ot. A második legnagyobb csoportot a Hadtudományi és Honvédtisztképző

Kar hallgatói alkotják, akik a teljes minta 17,7%-át és az NKE-s válaszadók 25,7%-át képviselik. A Rendészettudományi Kar hallgatói a válaszadók 10,7%-át adják, amely az NKE-s minta 15,5%-át jelenti. Végül a Víz tudományi Kar hallgatóit csupán egyetlen fő képviselte, ami a teljes mintában 0,5%-os arálynak felel meg, az NKE-s válaszadók körében pedig 0,7%-ot jelent.

Tovább vizsgálva a válaszokat, az 1.7. számú ábrán tüntettük fel, milyen típusú képzéseken vesznek részt a hallgatók. Ez esetben szintén több képzési formát jelölhettek a válaszadók. Akár csak a foglalkoztatási csoportoknál, itt sem fedeztünk fel anomáliákat a válaszok tekintetében.

A válaszadók képzési formák szerinti megoszlásának elemzése alapján az látható, hogy a legnagyobb arányt a mesterképzésben részt vevő hallgatók teszik ki, akik a teljes mintának 35,3%-át, összesen 76 főt képviselnek. Az alapképzés hallgatói a minta 32,6%-át teszik ki, amely 70 főt jelent. A doktori tanulmányokat folytató hallgatók a minta 3,7%-át alkotják, azaz mindössze 8 főt, míg a szakirányú továbbképzés hallgatói a válaszadók mindössze 1,4%-át képviselik, ami csupán 3 főt jelent. Az eredmények azt is mutatják, hogy az NKE hallgatói közül többen párhuzamos tanulmányokat



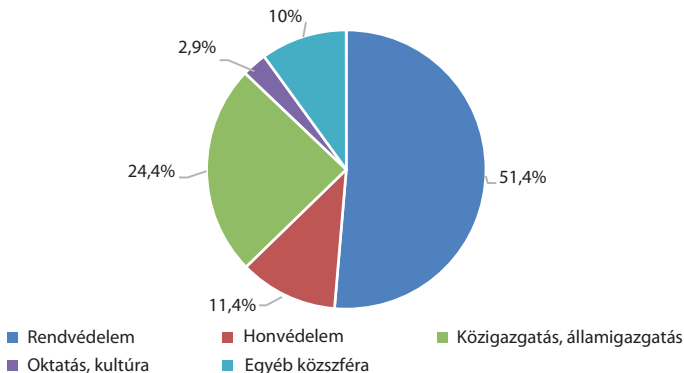
1.7. ábra: A Nemzeti Közszolgálati Egyetem hallgatóinak karonkénti megoszlása ($N = 148$)

Forrás: saját szerkesztés

folytatnak, hiszen a 148 NKE-jogviszonnyal rendelkező hallgató összesen 157 képzésben vesz részt. A keresztátlával végzett vizsgálat eredményei azt mutatták, hogy 7 hallgató folytat egyszerre alap- és mesterszakos tanulmányokat, 1 hallgató alapképzés mellett doktori képzésen vesz részt, 1 hallgató alapképzéses tanulmányai mellett szakirányú továbbképzésen is részt vesz, míg 1 mesterszakos hallgató doktori tanulmányokat is folytat.

A demográfiai mutatók vizsgálata során rákérdeztünk arra is, hogy a kitöltők a közsférában dolgoznak-e. A 215 válaszadó fele, 108 személy jelölte, hogy közsféra valamely területén dolgozik, aminek a megoszlását az 1.8. számú ábrán tüntettük fel.

A válaszadók közsféra szektorok szerinti megoszlása alapján a rendvédelem területéről érkező hallgatók alkotják a legnagyobb csoportot, összesen 36 főt, ami a teljes mintának 16,7%-át, a foglalkoztatási kategórián belül pedig 51,4%-ot tesz ki. A honvédelem területéről érkező válaszadók 8 főt képviselnek, ami a teljes mintának 3,7%-át és a közsférában dolgozók között 11,4%-os arányt jelent. A közigazgatás, államigazgatás szektorból érkező hallgatók száma 17 fő, ami a teljes minta 7,9%-át és a közsférás válaszadók



1.8. ábra: A közsférában dolgozók foglalkozási csoport szerinti megoszlása (N = 108)

Forrás: saját szerkesztés

24,4%-át teszi ki. Az oktatás, kultúra kategória csupán 2 főt számlál, ami a teljes minta 0,9%-át, a közszférás válaszadók között pedig 2,9%-ot tesz ki. Az egyéb közszféra kategóriába tartozók száma 7 fő, amely a teljes mintában 3,3%-ot, a közszférás válaszadók körében pedig 10%-ot képvisel.

Végezetül arra voltunk kíváncsiak, hogy a válaszadók korábban hallgattak-e bármilyen kiberbiztonsággal, adatvédelemmel kapcsolatos kurzust. A kérdőívben nem vizsgáltuk azt, hogy ez milyen időintervallumban, milyen képzési formában, típusban történt, vagy jelenleg részt vesz-e ilyen jellegű képzésben. Ez esetben a 215 válaszadó kevesebb mint fele, 94 kitöltő jelölte, hogy korábban már vett részt hasonló kurzuson.

2. VÉDEKEZÉS AZ ONLINE VISSZAÉLÉSEKKEL SZEMBEN

Az online visszaélések elleni védekezés az elmúlt években egyre nagyobb kihívássá vált köszönhetően az internetes térben megjelenő újabb és fejlettebb fenyegetéseknek. A digitális világ fejlődése nemcsak lehetőséget ad az embereknek a gyors kapcsolattartásra és információszerzésre, hanem teret nyit a kiberbűnözés különféle formái előtt is, mint például az adathalászat vagy a zsarolóvírusok. Az online visszaélések alatt minden olyan cselekményt, tevékenységet érthetünk, amely az internet vagy a digitális technológia felhasználásával történik. Ezek a visszaélések, csalások széles körben elterjedt problémává váltak a különböző területeken, a csalók folyamatosan új módszereket találnak a felderítő rendszerek gyengeségeinek kihasználására. Akár adathalászat, zsarolóvírusok vagy más módszerek révén, a digitális vagy online elemmel rendelkező visszaéléseket egyre nehezebb megállítani. Míg számos technikát fejlesztettek ki és alkalmaznak sikeresen a védelem oldalán, a támadók gyorsan alkalmazkodnak, és megtalálják a módját, hogy kijátsszák ezeket a rendszereket.¹³

TECHNOLÓGIAI FEJLŐDÉS ÉS A MESTERSÉGES INTELLIGENCIA

Az elmúlt néhány évtizedben óriási előrelépés történt a kommunikációs technológia fejlődésében, nevezetesen az interneten. Ez radikális változást hozott a tekintetben, hogy mit, milyen könnyen, mennyire kifinomultan és milyen minőségben lehet tenni a világhálón, valamint abban, hogy mindez világszerte elérhetővé válhatott a fogyasztók számára. Emellett rendkívül releváns a technológia fejlődése, ami többek között a képek (álló- és mozgóképek)

¹³ CHY 2024.

létrehozására és módosítására is összpontosít. A kommunikáció szempontjából az emberek erősen hajlamosak hinni a saját szemüknek és fülüknek, így a hang- és képi bizonyítékoknak hagyományosan nagy hitelt adnak a fényképeszeti trükkök hosszú története ellenére. Nem új az az ötlet, hogy egy képen egy arcot egy másik arcra cseréljenek. Például bárki, aki távolról is ismeri a grafikai tervezést, tanúsíthatja, hogy a különböző programok, mint például az Adobe Photoshop, viszonylag könnyen módosíthatják a digitális képeket. Ennek leküzdésére már régóta bevált technikák léteznek, amelyeket a képek hagyományos hamisításának és manipulációjának felderítésére használnak. Azonban a mesterséges intelligencia (*artificial intelligence*, a továbbiakban: AI), a mélytanulás és a képfeldolgozás közelmúltbeli térhódítása a *deepfake*¹⁴ videók előállításához vezetett, olyannyira, hogy nagyon könnyű valódi kinézetű mélyhamisított videót készíteni, és nehéz megkülönböztetni a valódi és a hamisított médiát. Az AI olyan számítógépes rendszerek elmélete és fejlesztése, amelyek képesek olyan feladatok elvégzésére, amelyek általában emberi intelligenciát igényelnek. Továbbá gyűjtőfogalom, amely az olyan komplex emberi képességek számításal történő reprodukálására irányuló erőfeszítések széles körét foglalja magában, mint a nyelvhasználat, a látás és az autonóm cselekvés. Az AI jelenlegi megközelítéseinek többsége gépi tanulási (*machine learning*, a továbbiakban: ML) technikákon alapul.¹⁵ Az ML-technikák arról ismertek, hogy képesek alkalmazkodni és tanulni. Bár az ML képes felismerni, megfékezni és megghiúsítani az ismert rosszindulatú támadásokat, a támadások bizonyos típusai túlmutatnak a jelenlegi kiberbiztonsági megoldásokon. Fontos még kiemelni, hogy az ML az AI egyik ága, amely statisztikai műveletek sorát használja a szükséges adatok kinyerésére és elemzésére, új jellemzők azonosítására, és hozzájárul a döntéshozatalhoz. Elsődleges célja, hogy a számítógép a szakemberek által bevitt adatokból

¹⁴ A *deepfake* olyan technológia, amely mesterséges intelligencia segítségével képes valóságként utánózni az emberek arcvonásait, hangját és mozdulatait, ezáltal szinte tökéletes másolatokat készítve. NMHH 2024. Másképpen: a *deepfake* digitálismédia-manipulációt jelent, olyan ultrarealisztikus, gépi tanúlással létrehozott hamis (ított) videót, amelynek a szereplői olyan dolgokat tehetnek vagy mondhatnak, amit a valóságban nagy valószínűséggel nem tennének vagy mondanának (VESZELSKI 2021, 2023).

¹⁵ CROSS 2022.

tanuljon, valamint számos olyan szabályt és módszert is magukban foglalnak, amelyek új adatmintákat vagy gyakorlatokat fedeznek fel vagy jósolnak meg. Ezek a technikák a kiberbiztonságban is hasznosíthatók, és felügyelt és felügyelet nélküli technikákba csoportosíthatók. Ezzel szemben a mélytanulási (*deep learning*, a továbbiakban: DL) modell kifejlesztésének két szakasza van. Az első lépés a terület titkosítása, amelyen az adatok a helyi környezetben a szerverre kerülnek. A második lépés ezen adatok elküldése a szerverre, és a szerverre érkező titkosított adatok feldolgozása, osztályozása és típusuk meghatározása. Például a karakterek képekről történő felismerése esetén az első szakasz a karakterek kódolását és továbbítását, míg a második szakasz a kapott adatok feldolgozását és a szerver és a helyi rendszer közötti *man-in-the-middle* támadás jelenlétének ellenőrzését foglalja magában, ami az adatok osztályozása szempontjából fontos. Ily módon biztosított a felhasználók számára az információk biztonságos továbbítása, és az, hogy illetéktelenek nem tudnak belépni a rendszerbe.¹⁶

Ahogy azt már az előbbiekben is olvashattuk, az ML-technikák létfontosságú szerepet játszanak a kiberbiztonság számos alkalmazásában, úgy többek között a különböző támadások korai felismerésében és előrejelzésében, mint például a *spamek* beazonosítása, a csalások felderítése, a rosszindulatú programok felderítése, az adathalászat, a *darkweb* vagy *deepweb* oldalak és a behatolásdetekció. Az ML-technikák megoldást jelenthetnek a kiberbűnözés felderítéséhez szükséges, e hiánypótló technológiákban jártas szakemberek hiányára. Emellett erőteljes megközelítésekre van szükség az új generációs (automatizált és evolúciós) kibertámadások észleléséhez és az ellenük való reagáláshoz. A gépi tanulás az egyik lehetséges megoldás az ilyen támadások elleni gyors fellépésre, mivel az ML képes tanulni a tapasztalatokból és időben reagálni az újabb támadásokra.¹⁷

Az elmúlt években a kiberbiztonsági alkalmazások széles körére próbáltak AI-alapú megoldásokat tervezni részben azért, mert a szervezetek egyre inkább megértik az AI jelentőségét a kiberfenyegetések mérséklésében. A nemlineáris problémák modellezésére szolgáló AI-alapú megközelítések például jól teljesítenek a nemlineáris osztályozásban, ami a kiberfenyegetések

¹⁶ MIJWIL-SALEM-ISMAEEL 2023.

¹⁷ SHAUKAT et al. 2020.

osztályozásának megkönnyítésére is felhasználható. Az AI-alapú megoldások iránti érdeklődés részben a számítástechnikai képességek fejlődésének is köszönhető. A Stanford Egyetem *AI Index 2019* jelentése szerint például a nagy méretű képosztályozó rendszer felhő-infrastruktúrán történő betanításához szükséges idő a 2017 októberében mért körülbelül három órától 2019 júliusára körülbelül 88 másodpercre csökkent. Az AI-alapú megközelítések számítási teljesítménye a jelentések szerint szintén körülbelül háromhavonta megduplázódik, meghaladva a Moore-törvényt. Ezek a képességek felhasználhatók az AI-alapú kiberbiztonsági megoldások teljesítményének javítására. Az AI-alapú megoldásokra példák az MIT és a PatternEx által kifejlesztett megoldások:

- ♦ Darktrace: AI-t használ a vállalati ellenálló-képesség kiépítésére;
- ♦ DeepArmor: AI-alapú rendszer a támadások ellen;
- ♦ X by Invincea: mély tanulást használ a biztonsági fenyegetések megértésére és észlelésére;
- ♦ Cognigo DataSense: amely gépi tanulási algoritmusokat használ az érzékeny adatok megkülönböztetésére és védelmére a nem érzékeny adatoktól.

Ugyanakkor az is ismert, hogy a gépi intelligencia nem helyettesítheti teljesen az emberi intelligenciát, és a mesterséges intelligencia következő generációja valószínűleg az emberi és a gépi intelligenciát fogja ötvözni (ezt más néven *human-in-the-loopnak* nevezik).¹⁸

A HUMÁN TÉNYEZŐ SZEREPE

A kiberbiztonsággal kapcsolatos kutatások többsége a hálózati rendszerek fejlesztésére összpontosít, mivel sokan úgy vélik, hogy az információtechnológiai fejlődés és a szoftverfejlesztés a fő módja az információbiztonság növelésének, azonban a támadók nem magát a számítógépes rendszert, hanem a számítógépes rendszer felhasználóit közvetlenül is manipulálhatják.

¹⁸ ZHANG et al. 2022.

Ezt megtehetik *social engineering* (például a számítógépes rendszer felhasználóinak becsapása információk, jelszavak megszerzése érdekében) és *kognitív hacking* (például téves információk terjesztése) segítségével, hogy betörjenek egy hálózatba vagy számítógépes rendszerbe. Moustafa, Bello és Maurushat, a Nyugat-Sydney-i Egyetem kutatóinak 2021-es tanulmánya szerint a *social engineering* támadások az összes kiberbiztonsági támadás 28%-át teszik ki, és 24%-uk adathalászat miatt történt, illetve a *social engineering* támadások több mint 70%-a sikeres volt az elmúlt években, ezáltal elmondható, hogy az emberi hibák jelentik az egyik legnagyobb fenyegetést a kiberbiztonságban. A statisztikai adatok szerint tehát az adathalászt támadások voltak a leggyakoribb támadások, és részben *social engineeringet* használtak fel arra, hogy az áldozatokat rosszindulatú szoftverek telepítésére vagy rosszindulatú webhelyek látogatására rávegyék, így megszerezve a hitelesítő adataikat. Az ilyen típusú támadások során az áldozatoknak gyakran küldenek olyan e-maileket vagy szöveges üzeneteket, amelyek például úgy tűnnek, mintha egy szoftverfrissítésről, egy harmadik féltől származó, legitim levelezésről, egy aktuális katasztrófával vagy válsággal kapcsolatos információról vagy egy bank, esetleg éppenséggel egy közösségi oldal értesítéséről szólnának. Az adathalász támadásoknak való áldozatul esésén kívül a számítógépes rendszerek felhasználói más kiberbiztonsági hibákat is elkövetnek, például megosztják a jelszavakat barátaikkal és családtagjaikkal, vagy nem telepítik a szoftverfrissítéseket.¹⁹ Az empirikus eredmények alapján a számítógépes alapismeretek és az információkeresési készségek is befolyásolhatják az egyén viselkedését a biztonság kezelésében. Emellett hangsúlyos a tudásmegosztás fontossága is a fenyegetések, valamint az információbiztonság költségeinek csökkentése érdekében. Ebből következik, hogy az egyének – üzleti környezetben a munkavállalók – kiberbiztonsággal kapcsolatos ismeretei jelentős hatással vannak a kiberbiztonsági tudatossági képességek kialakítására. A biztonsági ismeretek szintjét úgy határozzák meg, hogy az egyén mennyire ismeri a kiberfenyegetésekkel, sebezhetőségekkel, támadási mintákkal és ezek rendszerre gyakorolt hatásával kapcsolatos elméleti információkat. Számos társadalmi, pszichológiai és demográfiai

¹⁹ MOUSTAFA–BELLO–MAURUSHAT 2021.

tényező, mint például az életkor, a nem stb. befolyásolhatja az alkalmazottak viselkedését a kiberfenyegetések kezelésében. Korábban a védelmi motiváció elméletét alkalmazták annak magyarázatára, hogy a különböző nemű munkavállalók miért viselkednek eltérően a kiberfenyegetések kezelésében. A legújabb tanulmányok azonban kimutatták, hogy a biztonság tudatos viselkedés összefügg az olyan változókkal, mint az észlelt érzékenység, az észlelt súlyosság, az észlelt előnyök és az önhatékonyság. Az egyének kiberfenyegetéssel kapcsolatos hozzáállását gyakran figyelmen kívül hagyják, mivel nehéz felmérni a tudatossági szintjét és azt, hogy rendelkezik-e megfelelő készségekkel a kiberfenyegetésekkel szembeni védekezéshez.²⁰

VÉDEKEZÉS AZ ONLINE TÉRBEN

Az online térben való védekezés napjaink egyik legfontosabb biztonsági kihívása, hiszen a digitális eszközök és hálózatok használata mindennapi életünk része lett. Az internet nyújtotta előnyök mellett a kiberbiztonsági fenyegetések – például a korábban már említett adathalászat, a zsaroló-programok és a *deepfake* – egyre kifinomultabb módszerekkel próbálnak hozzáférni érzékeny információinkhoz. A megfelelő védekezési módszerek alkalmazása ezért kulcsfontosságúvá vált a személyes adatok és rendszerek védelme érdekében.

Különböző források különböző jógyakorlatokat említenek a kiberbiztonságos viselkedés kapcsán, több témakörre osztva, ezek a következők.

Jelszavak és többfaktoros hitelesítés

A jelszavak állnak az adataink és azok között, akik ki akarják használni azokat. Ezért kritikus fontosságú, hogy olyan összetett jelszavakat dolgozzunk ki, amelyeket szinte lehetetlen kitalálni. Hogyan is néz ki egy erős jelszó? Az erős jelszavak legalább 12, de lehetőleg több mint 16 karakter hosszúak

²⁰ AKTER et al. 2022.

és kis- és nagybetűket, valamint különleges karaktereket is tartalmaznak (#, ! stb.). Lehetőség szerint kerülnünk kell a személyes adatok, például a nevek vagy a születési dátumok használatát. Amennyiben egy jelszónak megjegyezhetőnek kell lennie, próbáljunk meg egy szó helyett egy kifejezést használni. Az összetett jelszavak megléte csak az első lépés az adatok védelmében. Célszerű bevezetni egy olyan jelszókezelő eszközt is, amelyet a hitelesítő adatok létrehozására, tárolására, megosztására és kezelésére használnak. Emellett tartózkodnunk kell a titkosítatlan jelszavak e-mailben, sms-ben vagy más módszerekkel történő megosztásától. Az így küldött jelszavak könnyen kompromittálhatók. Az e-mail-fiókokat például viszonylag könnyű feltörni, és a kiberbűnözőkről ismert, hogy a *man-in-the-middle* támadással sikeresen elfogják az e-maileket és az üzeneteket.

Emellett a közösségi profilokon, valamint a banki rendszerek esetén célszerű a többfaktoros hitelesítés (*multi-factor authentication*, a továbbiakban: MFA) engedélyezése, amely megköveteli, hogy két vagy több módszerrel igazoljuk személyazonosságunkat. Például a jelszó megadása után egy alkalmazás segítségével hitelesítenünk kell magunkat, és meg kell adni egy egyszeri jelszót, amelyet szöveges üzenetben küldenek az eszközünkre. Az MFA használatával még ha egy harmadik fél hozzá is jut a jelszóhoz, a többi lépés megtétele nélkül nem fér az alkalmazáshoz vagy szolgáltatáshoz. A Microsoft szerint az MFA bevezetésével a fiókok elleni támadások 99,9%-a megelőzhető.

Rendszeres szoftverfrissítés

Amikor a teendőink listája kilométer hosszúságú, könnyű a „később emlékeztetni” gombra kattintani, amikor megjelenik a „készülékének frissítésre van szüksége” felugró ablak. Azonban nem szabad figyelmen kívül hagyni ezeket a frissítéseket. Sok eszközfrissítés biztonsági javításokat tartalmaz, amelyek kritikus fontosságúak az eszköz védelme szempontjából. Akár üzleti, akár privát használatú eszközről van szó, a frissítéseket mindig telepíteni kell, amikor azok elérhetők.

Biztonsági szoftverek használata

Ennek egyik aspektusa a távoli munkavégzés során elterjedt virtuális magán-hálózat (*virtual private network*, a továbbiakban: VPN) használata. A VPN biztonságos és titkosított kapcsolatot hoz létre az eszközünk és a csatlakoztatott hálózat között, amely kiemelt fontosságú, amikor üzleti alkalmazásokat és szolgáltatásokat érünk el. A VPN titkosítja az eszköz és a hálózat között megosztott üzleti adatokat, így azok (még nyilvános internetkapcsolat használata esetén is) titokban maradnak. Ennek eredményeképpen az üzleti adatok védve vannak a kíváncsi szemektől.

Emellett minden internethez csatlakozni tudó eszköz egy potenciális ajtó, amelyet a támadók arra használhatnak, hogy hozzáférjenek az adatainkhoz – itt jön szóba a tűzfal és a vírusirtó használata. A vírusirtó eszközök megakadályozzák, hogy a rosszindulatú szoftverek megfertőzzék az eszközt, míg a tűzfalak a hálózaton bejövő és kimenő forgalom felügyeletével védik az adatokat. A beállított biztonsági szabályok alapján a tűzfal engedélyezheti vagy megtagadhatja a forgalmat.

Tudatos viselkedés

A kiberbűnözők által alkalmazott egyik leggyakoribb módszer az adathalászat. Ez akkor történik, amikor egy csaló olyan e-mailt vagy szöveges üzenetet küld, amely arra kéri, hogy információkat adjunk meg. És sajnos egyre jobbak, hitelesebbnek tűnnek ezek az üzenetek. Például kaphatunk olyan e-mailt, amelyben azt írják, hogy a hitelesítő adataink veszélybe kerültek, és arra ösztönöznek minket, hogy egy hamis weboldalon jelentkezünk be. Amint ez megtörténik, a támadó már hozzáférést szerez a fiókunkhoz. A tudatos viselkedés kialakításához elsősorban azt szükséges tudnunk, hogyan néznek ki ezek az adathalász e-mailek. Jellemzően sok helyesírási hibát tartalmazhatnak, és nem szólítják meg név szerint a személyt – a magyar nyelv sajátossága miatt gyakran keverik a tegezést, illetve a magázást egy üzeneten belül. Emellett lehet, hogy tartalmaznak egy kattintható linket vagy egy letölthető mellékletet is. Ennek érdekében tartózkodnunk kell a gyanúsnak tűnő e-mailek (vagy szöveges üzenetek)

megnyitásától és az ismeretlen linkekre való kattintástól. Fontos azonban észben tartanunk, hogy a mesterséges intelligencia segítségével a támadók egyre kifinomultabb és nehezen azonosítható kampányokat hajthatnak végre.

Mindezek mellett célszerű bizonyos időközönként ellenőrizni online fiókjaink esetében, hogy nincs-e gyanús tevékenységről vagy jogosulatlan módosításról szóló értesítés. Továbbá kiemelten fontos, hogy naprakészen tartsuk a kiberbiztonsággal kapcsolatos ismereteinket, felismerve, hogy a kiberbiztonság állandó folyamat. Legyünk éberek a kiberfenyegetésekkel szemben!²¹

MÓDSZER

A közösségimédia-platformokon és tágabb értelemben az online felületeken és a digitális eszközök használatában mutatkozó aktív védekező viselkedések mérésére kérdőívünkben külön szakasz szolgált. Általános értelemben aktív védekező viselkedésnek tekintettünk minden olyan stratégiai (vagyis tudatos és szándékos) viselkedést, amelynek célja az online visszaélések megelőzése, illetve megfelelő kezelése. E kérdőívblokk tételeinek kidolgozásakor az úgynevezett induktív szerkesztési módot alkalmaztuk.²² Kiemelt szempont volt a tartalmi érvényesség kritériumának teljesítése, vagyis hogy a tételek együttesen lefedjék azon viselkedések valamennyi fő típusát, amelyek a modern online környezetben hozzájárulhatnak a felhasználó biztonságához. A fejlesztés első fázisában nagyszámú tételt konstruáltunk egyrészt a fent hivatkozott források alapján, másrészt saját kiberbiztonsági képzéseink tananyagaiból merítve. A második szakaszban rendszerezettük a tételeket, megszüntettük a tartalmi átfedéseket az itemek átfogalmazásával és összevonásával, továbbá kiválogattuk azokat a viselkedéseket, amelyek esetében várhatóan változatosan alakul a felhasználói magatartás. Végül a tételeket csoportosítva kialakítottunk egy struktúrát, amely a védekező viselkedéseket két fődimenzióra, azokon belül pedig két-két aldimenzióra bontja.

Az így megszerkesztett modellünkben az A fődimenzióhoz tartoznak mindazon viselkedések, amelyek célja az információ hozzáféréseinek vagy

²¹ Coursera Staff 2024; BAIG 2024; TeamPassword 2024.

²² SZOKOLSZKY 2004.

az információ terjedése hatékony ellenőrzésének, kontrollálásának biztosítása. E fődimenzió alatt két aldimenzió különböztethető meg az információ terjesztése szerint. Az A₁ aldimenzióhoz soroltunk minden olyan viselkedésmódot, amellyel a felhasználó a személyes adataihoz való hozzáférést igyekszik megakadályozni: ezek tehát azt a célt szolgálják, hogy illetéktelen személyek ne férjenek hozzá olyan információkhoz, amelyeket a felhasználó nem tett közzé. E viselkedések körén belül is elkülöníthetők specifikusabb kategóriák: például némelyek azt szolgálják, hogy a felhasználó megelőzze a személyes adatai (jelszavak, banki adatok, lakcím, e-mail-cím stb.) illetéktelen személyekhez kerülését, mások pedig arra irányulnak, hogy a felhasználó hatékonyan elhárítsa az esetleges negatív következményeket, ha ilyen eset mégis előfordulna. Elemzésünkben azonban az A₁ aldimenziót nem bontottuk tovább specifikusabb kategóriákra: mint később látni fogjuk, hierarchikus modellünk így is négy különböző szinten tette lehetővé az összefüggések vizsgálatát: (1) az aktív védekező viselkedések szintjén általános értelemben, (2) az A és B fődimenziók szintjén, (3) az A₁, A₂, B₁, B₂ aldimenziók szintjén és (4) a kérdőívtelekben megfogalmazott konkrét viselkedések szintjén. Az A₁ aldimenzióhoz tartozik tehát minden olyan viselkedés, amelynek az a célja, hogy a felhasználó megakadályozza illetéktelen személyek hozzáférését olyan személyes információkhoz, amelyeket másokkal nem osztott meg és nem is kíván megosztani. E kategóriába tartozó kérdőívteleinket a 2.1. táblázat tartalmazza.

Az A₁₀₇R kérdés kivételes abban az értelemben, hogy úgynevezett fordított tétel. E tétel esetében ugyanis a megfogalmazott viselkedés óvatlanságot fejez ki, és a kitöltő egyet nem értése felel meg annak, hogy odafigyel az újonnan telepített alkalmazások hozzáférési jogosultságaihoz, vagyis igyekszik védekezni az ellen, hogy különféle alkalmazások a személyes adatait összegyűjtsék. E tétel esetében azért döntöttünk a fordított megfogalmazás mellett, hogy kerüljük a tagadó szerkezet használatát („nem engedélyezem a szükségtelennek látszó hozzáféréseket”), mivel – különösen terjedelmesebb kérdőívek esetében – gyakran tapasztalható, hogy a tagadással való egyetértés/egyet nem értést a kitöltők tévesen értelmezik. Az összesített mutató számításakor e tétel fordított jellegét természetesen figyelembe vettük, és a tételre adott válasz pontértékét tükröztük. A tételek szintjén végzett elemzésekben ugyanakkor a kitöltő válaszában eredeti értékével számoltunk.

2.1. táblázat: Az aktív védekező viselkedési formák A1 (személyes adatokhoz való hozzáférés ellenőrzése) aldimenziójához tartozó kérdőívitételek

Azonosító	Szöveg
A101	Rendszeresen frissítem a jelszavaimat a közösségimédia-fiókjaimban.
A102	Minden platformhoz más jelszót használok.
A103	Odafigyelek arra, hogy a jelszavaim ne tartalmazzanak velem kapcsolatba hozható információt (mint például nevek, monogramok, születési dátumok, vagy hobbi, tulajdonra utaló szavak).
A104	Jelszókezelő alkalmazást használok.
A105	Rendszeresen ellenőrzöm a fiókjaim hozzáférési naplóit, hogy azonosítsam az esetleges illetéktelen hozzáféréseket.
A106	Kétfaktoros hitelesítést használok a közösségimédia-fiókjaim védelme érdekében (a fiókhoz való hozzáféréshez két lépés szükséges, például a jelszavam megadása, majd a mobiltelefonomra érkező egyszeri kód megadása).
A107R	Az újonnan telepített mobilalkalmazásoknak engedélyezem a kért hozzáféréseket (például a kamerához, a tártózkodási helyhez, a kapcsolatokhoz stb.).
A108	Rendszeresen ellenőrzöm a mobilkészülékemre telepített alkalmazások hozzáférési jogosultságait.
A109	Ha gyanús tevékenységet észlelek a közösségimédia-fiókomban, azonnal megváltoztatom a jelszavam.
A110	Ha gyanús tevékenységet észlelek a közösségimédia-fiókomban, azonnal értesítem a platform ügyfélszolgálatát.

Forrás: saját szerkesztés

Az A fődimenzió (hozzáférés ellenőrzése) belül különálló kategóriát alkottak azok a viselkedések, amelyek a felhasználó által tudatosan közzétett tartalmakra vonatkoznak, és amelyek célja a közzétett tartalmak körének korlátozása, illetve azok terjedésének kontrollálása. Modellünkben ezek tartoznak az A2 aldimenzióba, és a kérdőívünkben a 2.2. táblázatban felsorolt tételek feleltek meg e kategóriának.

2.2. táblázat: Az aktív védekező viselkedések A₂ (közzétett tartalmakhoz való hozzáférés ellenőrzése) aldimenziójához tartozó kérdőívtelemek

Azonosító	Szöveg
A201	Tudatosan kerülöm a személyes adataim (például e-mail-cím, telefonszám, lakcím, tartózkodási hely) online megosztását.
A202	Aktívan figyelem, hogy a közösségi médiában megjelenő személyes információim kiszivárogtak-e.
A203	A közösségi média használatakor előnyben részesítem a privát kommunikációs csatornákat (például üzenetküldés, privát csoportok).
A204	A közösségimédia-profiljaimat privátra állítom.
A205	Nem hivatalos alkalmazásboltból származó alkalmazásoknak csak indokolt esetben engedélyezem a hozzáférést a fiókom adataihoz.
A206	Nem járulok hozzá a fiókadataim marketingcélú felhasználásához (például személyre szabott hirdetések küldéséhez).
A207	Rendszeresen ellenőrzöm az általam használt közösségi portálok tevékenységi naplóját, és törölöm azokat a tevékenységeket, amelyek érzékeny, személyes információkkal kapcsolatosak.
A208	A megfelelő beállításokkal korlátozom, kik láthatják a közösségi médiában a profilinformációimat.
A209	A megfelelő beállításokkal korlátozom, kik kaphatnak hozzáférést a közösségi médiában a posztjaimhoz.
A210	A megfelelő beállításokkal korlátozom, kik kereshetnek rám e-mail-cím vagy telefonszám alapján.
A211	A megfelelő beállításokkal korlátozom a profilom láthatóságát keresőmotorokban.
A222	Rendszeresen ellenőrzöm a közösségimédia-fiókjaim követőinek listáját.

Forrás: saját szerkesztés

Az A fődimenzióba tartozó viselkedésektől határozottan elkülönülnek azok, amelyek azt a célt szolgálják, hogy a felhasználó ne maradjon le az online fenyegetések világában: ezek alkotják a B fődimenziót, amelyet „naprakészségnek” neveztünk el. E dimenzión belül is két aldimenzió különíthető el aszerint, hogy a naprakészség mire vonatkozik. A B aldimenzióba tartoztak azok a viselkedések, amelyek azt a célt szolgálják, hogy a felhasználó

megfelelően tájékozódjon arról, milyen új fenyegetések jelenthetnek veszélyt az online térben, és milyen eszközök állnak rendelkezésre az ellenük való védekezésre. A B1 aldimenzióba tartozó viselkedések tehát mind az ismeretek naprakészen tartását szolgálják. Ha finomabb felosztású elméleti struktúra megalkotására törekednénk, természetesen ezek is tovább bonthatók lennének különféle alkategóriákra, például a tájékoztatói források és csatornák taxonómiája alapján. Jelen elemzésünkben azonban nem törekedtünk ennél további, köztes rendezőelv kialakítására az aldimenziók és a tételek szintje között. A B1 aldimenzióhoz kérdőívünkben a 2.3. táblázatban látható tételek tartoztak.

2.3. táblázat: Az aktív védekező viselkedések B1 (tájékoztató) aldimenziójához tartozó kérdőív-tételek

Azonosító	Szöveg
B101	Rendszeresen tájékozódok a legújabb online biztonsági fenyegetésekről és azok elhárításáról.
B102	Elolvatom a közösségimédia-alkalmazások felhasználói szerződéseit mielőtt telepíteném őket.
B103	Elolvatom a közösségimédia-platformok adatvédelmi irányelveit mielőtt használnám őket.
B104	Elolvatom a közösségimédia-platformoktól érkező, biztonsági frissítésekkel kapcsolatos tájékoztató üzeneteket.
B105	Aktívan keresek a közösségimédia-biztonsággal kapcsolatos online tartalmakat.
B106	Az általam követett hírportálokon elolvasom az online biztonsági fenyegetésekkel kapcsolatos cikkeket.
B107	Szívesen nézek olyan videókat, amelyek bemutatják, hogyan védekezhetünk az online biztonsági fenyegetések ellen.
B108	Követek olyan csatornát/blogot, amely kifejezetten az online biztonság kérdéseivel foglalkozik.
B109	Aktívan részt veszek online közösségekben, amelyek a digitális biztonsággal és adatvédelemmel foglalkoznak.
B110	Rendszeresen beszélgetek ismerőseimmel az online visszaélések veszélyeiről és a védekezés lehetőségeiről.

Forrás: saját szerkesztés

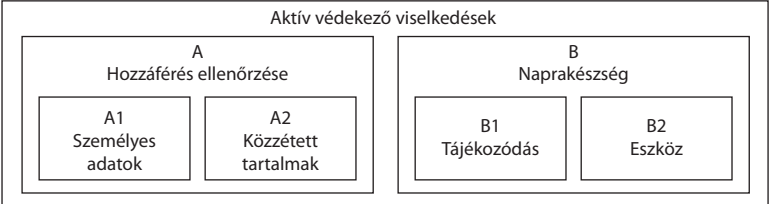
Az ismeretek naprakészen tartása mellett külön viselkedéskategóriát alkotnak a védekezés azon módjai, amelyek a felhasználó hardvereszközeinek naprakészen tartását szolgálják. Vizsgálatunkban e viselkedés alkotta a B2 aldimenziót, amelyet a 2.4. táblázatban felsorolt öt kérdőívtételén keresztül mértünk.

2.4. táblázat: Az aktív védekező viselkedések B2 (eszközök naprakészen tartása) aldimenziójához tartozó kérdőívtételek

Azonosító	Szöveg
B201	Odafigyelek, hogy a közösségimédia-alkalmazásaimban be legyen kapcsolva az automatikus frissítés funkciója.
B202	Ha értesítést kapok arról, hogy elérhető a mobileszközöm operációs rendszerének (iOS, Android) frissítése, igyekszem minél hamarabb telepíteni az új verziót.
B203	Ha egy alkalmazás automatikus frissítéséről kapok értesítést, igyekszem minél hamarabb telepíteni azt.
B204	Rendszeresen ellenőrzöm, elérhetők-e frissítések az eszközeimre telepített alkalmazásokhoz (az automatikus frissítéseken túlmenően).
B205	Ha egy alkalmazást a legújabb verzióra frissítettem, ellenőrzöm, változtak-e az adatvédelmi beállítások lehetőségei.

Forrás: saját szerkesztés

Az aktív védekező viselkedések általunk alkalmazott hierarchikus modelljét a 2.1. ábra összegzi.



2.1. ábra: Az aktív védekező viselkedések hierarchikus modellje

Forrás: saját szerkesztés

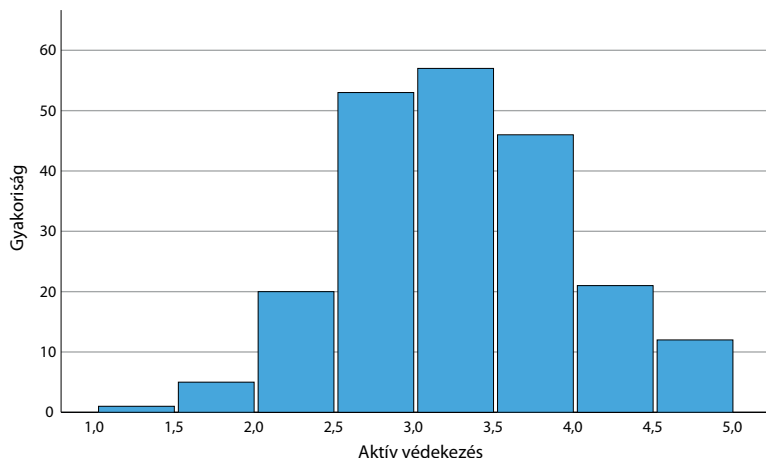
E kérdőívblokkban a kitöltők valamennyi viselkedésről oly módon nyilatkoztak, hogy egy Likert-mátrixban jelölték, milyen mértékben jellemzők rájuk az egyes tételekben megfogalmazott állítások. A blokk elején olvasható instrukció a következő volt:

Ön konkrétan mit tesz azért, hogy megelőzze vagy elhárítsa az online visszaélések veszélyét? Milyen mértékben jellemzők Önre az alábbi viselkedések?

- 1 – Egyáltalán nem jellemző rám.
- 2 – Inkább nem jellemző rám.
- 3 – Változó: nem tudom eldönteni.
- 4 – Inkább jellemző rám.
- 5 – Kifejezetten jellemző rám.

Az egyes fő- és aldimenzióknak megfelelő összesített mutatókat az adott kategóriába tartozó tételekre adott válaszok átlaga adta. A tételstruktúra elemzése alapján a mérőeszköz minden szinten jó vagy kiváló megbízhatósággal mérte a vizsgált jelenséget. Az A₁, A₂, B₁, B₂ aldimenziók Cronbach-alfa mutatója rendre 0,76; 0,88; 0,93; illetve 0,84. Az A és B fődimenziók Cronbach-alfája 0,89, illetve 0,92, és az aktív védekező viselkedések összesített mutatójának megbízhatósága 0,94.

Az összesített mutatók eloszlása az elemzés valamennyi szintjén halom alakú, és közelítőleg szimmetrikus eloszlást mutat. A 2.2. ábrán látható hisztogram a két fődimenziót összesítő általános mutató eloszlását szemlélteti.



2.2. ábra: Az aktív védekező viselkedések összesített mutatójának eloszlása

Forrás: saját szerkesztés

EREDMÉNYEK

A leíró statisztikák azt mutatják, hogy az aktív védekező viselkedések mérésére szolgáló tételek egyikénél sem kaptunk padló- vagy plafoneffektust. A legelterjedtebb és legritkábban megjelenő viselkedések esetében a szórás kis mértékben alacsonyabb volt, mint más tételeknél, és ezekben az esetekben az eloszlás ferdeséget is mutatott, de a terjedelem valamennyi tétel esetében kitöltötte az ötfokú skálát.

Ha a válaszokat tételenként összesítjük, megmutatkozik, mely viselkedések a leggyakoribbak és a legkevésbé jellemzőek kitöltőink körében. E sorrendet mutatja a 2.5. táblázat.

2.5. táblázat: Az aktív védekező viselkedést
mérő kérdőívtetelek leíró statisztikái

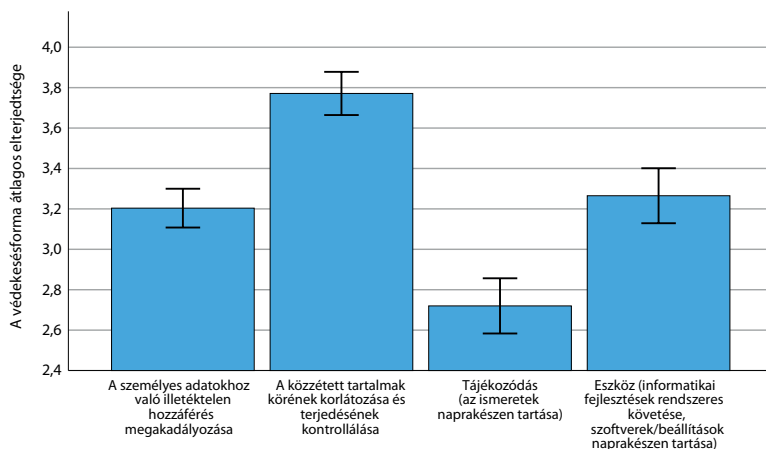
	Átlag	Szórás
A közösségi média használatakor előnyben részesítem a privát kommunikációs csatornákat (például üzenetküldés, privát csoportok).	4,25	0,932
Tudatosan kerülöm a személyes adataim (például e-mail-cím, telefonszám, lakcím, tartózkodási hely) online megosztását.	4,11	1,066
Nem járulok hozzá a fiókadataim marketingcélú felhasználásához (például személyre szabott hirdetések küldéséhez).	4,04	1,041
Ha gyanús tevékenységet észlelek a közösségimédia-fiókomban, azonnal megváltoztatom a jelszavam.	3,99	1,066
A megfelelő beállításokkal korlátozom, kik kaphatnak hozzáférést a közösségi médiában a posztjaimhoz.	3,97	1,123
A megfelelő beállításokkal korlátozom, kik láthatják a közösségi médiában a profilinformációimat.	3,93	1,210
Nem hivatalos alkalmazásboltból származó alkalmazásoknak csak indokolt esetben engedélyezem a hozzáférést a fiókom adataihoz.	3,93	1,198
A közösségimédia-profiljaimat privátra állítom.	3,91	1,294
Ha értesítést kapok arról, hogy elérhető a mobilkészítőm operációs rendszerének (iOS, Android) frissítése, igyekszem minél hamarabb telepíteni az új verziót.	3,71	1,243
A megfelelő beállításokkal korlátozom, kik kereshetnek rám e-mail-cím vagy telefonszám alapján.	3,71	1,294
Kétfaktoros hitelesítést használok a közösségimédia-fiókjaim védelme érdekében (a fiókhoz való hozzáféréshez két lépés szükséges, például a jelszavam megadása, majd a mobiltelefonomra érkező egyszeri kód megadása).	3,65	1,262
Ha egy alkalmazás automatikus frissítéséről kapok értesítést, igyekszem minél hamarabb telepíteni azt.	3,55	1,198
Odafigyek arra, hogy a jelszavam ne tartalmazzanak velem kapcsolatba hozható információt (mint például nevek, monogramok, születési dátumok vagy hobbi, tulajdonra utaló szavak).	3,55	1,248
A megfelelő beállításokkal korlátozom a profilom láthatóságát keresőmotorokban.	3,53	1,349
Rendszeresen ellenőrzöm a közösségimédia-fiókjaim követőinek listáját.	3,43	1,287
Rendszeresen ellenőrzöm az általam használt közösségi portálok tevékenységi naplóját, és törölöm azokat a tevékenységeket, amelyek érzékeny, személyes információkkal kapcsolatosak.	3,30	1,307
Rendszeresen ellenőrzöm, elérhető-e frissítések az eszközeimre telepített alkalmazásokhoz (az automatikus frissítéseken túlmenően).	3,28	1,321
Minden platformhoz más jelszót használok.	3,19	1,198

	Átlag	Szórás
Szívesen nézek olyan videókat, amelyek bemutatják, hogyan védekezhünk az online biztonsági fenyegetések ellen.	3,15	1,336
Aktívan figyelem, hogy a közösségi médiában megjelenő személyes információim kiszivárogtak-e.	3,14	1,254
Rendszeresen ellenőrzöm a mobilszközömre telepített alkalmazások hozzáférési jogosultságait.	3,12	1,262
Odafigyelek, hogy a közösségimédia-alkalmazásaimban be legyen kapcsolva az automatikus frissítés funkciója.	3,12	1,308
Rendszeresen tájékozódok a legújabb online biztonsági fenyegetésekről és azok elhárításáról.	3,07	1,202
Az általam követett hírportálokon elolvasom az online biztonsági fenyegetésekkel kapcsolatos cikkeket.	3,07	1,338
Rendszeresen frissítem a jelszavaimat a közösségimédia-fiókjaimban.	3,01	1,207
Ha gyanús tevékenységet észlelek a közösségimédia-fiókomban, azonnal értesítem a platform ügyfélszolgálatát.	2,98	1,350
Rendszeresen beszélgetek ismerőseimmel az online visszaélések veszélyeiről és a védekezés lehetőségeiről.	2,97	1,315
Rendszeresen ellenőrzöm a fiókjaim hozzáférési naplóit, hogy azonosítsam az esetleges illetéktelen hozzáféréseket.	2,82	1,271
Az újonnan telepített mobilalkalmazásoknak engedélyezem a kért hozzáféréseket (például a kamerához, a tartózkodási helyhez, a kapcsolatokhoz stb.).	2,77	1,300
Követek olyan csatornát/blogot, amely kifejezetten az online biztonság kérdéseivel foglalkozik.	2,72	1,394
Ha egy alkalmazást a legújabb verzióra frissítettem, ellenőrzöm, változtak-e az adatvédelmi beállítások lehetőségei.	2,67	1,401
Elolvasom a közösségimédia-platformoktól érkező, biztonsági frissítésekkel kapcsolatos tájékoztató üzeneteket.	2,61	1,317
Aktívan keresek a közösségimédia-biztonsággal kapcsolatos online tartalmakat.	2,56	1,228
Jelszókezelő alkalmazást használok.	2,50	1,475
Aktívan részt veszek online közösségekben, amelyek a digitális biztonsággal és adatvédelemmel foglalkoznak.	2,37	1,271
Elolvasom a közösségimédia-platformok adatvédelmi irányelveit, mielőtt használnám őket.	2,35	1,273
Elolvasom a közösségimédia-alkalmazások felhasználói szerződéseit mielőtt telepíteném őket.	2,34	1,290

Forrás: saját szerkesztés

E leíró statisztikákból is szembetűnő, hogy a legelterjedtebb viselkedések nagy része az A2 aldimenzióhoz tartozik. Tehát az eredmények azt mutatják, hogy az adatközlők sokat tesznek azért, hogy kontroll alatt tartsák, kivel milyen információkat osztanak meg, odafigyelnek arra, hogy ezek az információk az ismerőseik szűk körén belül maradjanak, és se illetéktelen személyekhez, se gazdasági társaságokhoz ne jussanak el. Ha a legkevésbé jellemző viselkedéseket vesszük szemügyre, jól látható, hogy a legritkábban tanúsított a tájékozódás (B1 aldimenzió) kategóriájába tartoznak: a kitöltők ritkán veszik a fáradságot arra, hogy adatvédelmi irányelveket vagy a felhasználói szerződések pontos tartalmát megismerjék, de a kiberbiztonsággal kapcsolatos ismeretterjesztő cikkek vagy a témával foglalkozó csoportok iránt is csekély az érdeklődés.

A négy aldimenzió leíró statisztikai megerősítik ezt a mintázatot. A négy viselkedéskategória átlagos elterjedtségét a 2.3. ábra mutatja, 95%-os konfidenciaintervallumokkal.



2.3. ábra: Az aktív védekező viselkedések négy aldimenziójának elterjedtsége (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

Az egyes viselkedéskategóriák elterjedtsége közötti különbségeket a statisztikai próba eredménye is megerősíti. Az előzetes elemzéseink azt mutatták, hogy az adatsorban található mintázat szignifikánsan sérti a szfericitás feltételét, vagyis az egyes feltételpárok közötti különbségek varianciája nem egyenlő (Mauchly-féle próba: $W = 0,789$, $\chi^2(5) = 50,466$, $p < 0,001$). Ilyen esetben a szfericitás sérülése az ismételt mérés varianciaanalízis eredményeit torzíthatja, és bár a hatások korrigálására több eljárás is használható, célszerű alternatíva az elemzést multivariáns varianciaanalízissel (MANOVA) elvégezni, mivel ez utóbbi eljárás nem igényli a szfericitás teljesülését.²³ Az elemzés eredménye szerint a négy viselkedéstípus elterjedtsége közötti eltérések erősen szignifikánsak (Pillai-féle nyom: $V = 0,565$; $F(3, 212) = 91,901$; $p < 0,001$).

ÖSSZEGZÉS

A fejezet az online visszaélések elleni védekezés modern kihívásait és megoldásait vizsgálja, különös tekintettel az aktív védekező viselkedésekre. A bevezetésben kiemeltük, hogy a kiberbiztonsági fenyegetések, például az adathalászat vagy a zsarolóprogramok, folyamatosan fejlődnek, és a technológiai megoldások önmagukban nem elégségesek. Az egyéni védekező stratégiák és a tudatos viselkedés szerepe kulcsfontosságú a kiberbiztonsági rendszerek erősítésében.

Az eredmények alátámasztják a bevezetés állításait, miszerint az emberi tényező kiemelkedő szerepet játszik az online biztonságban. Az aktív védekező viselkedések vizsgálata négy fő kategóriára bontva történt:

- ♦ hozzáférés ellenőrzése (A1: személyes adatok és A2: közzétett tartalmak aldimenziók).
- ♦ naprakészség (B1 és B2 aldimenziók).

²³ FIELD 2009.

Az eredmények szerint a leggyakoribb védekező viselkedések az információhozzáférés kontrolljával kapcsolatosak:

- A személyes adatok online megosztásának elkerülése (4,11 átlagérték).
- Privát kommunikációs csatornák használata (4,25 átlagérték). Ezek azt mutatják, hogy a felhasználók elsősorban a közvetlen, azonnali veszélyek elkerülésére fókuszálnak.

Ugyanakkor a naprakészséghez kapcsolódó viselkedések, például az online biztonsági fenyegetésekről való rendszeres tájékozódás (3,07 átlagérték) és az adatvédelmi irányelvek elolvasása (2,35 átlagérték) jelentősen alacsonyabb gyakoriságot mutattak. Ez alátámasztja, hogy bár a felhasználók óvatosak bizonyos szempontból, kevésbé hajlandók időt és energiát fordítani a mélyebb információgyűjtésre.

A bevezetésben említett technológiai fejlődés, például a mesterséges intelligencia alkalmazása, fontos szerepet kapott az online biztonság növelésében. Az eredmények azonban arra utalnak, hogy az emberi tényezők hiányosságai – például a tájékozottság alacsony szintje – jelentős kihívásokat okoznak. Míg a technológiai eszközök fejlődése gyors, az egyéni viselkedésbeli hiányosságok továbbra is sebezhetőségeket teremtenek.

Elmondható tehát, hogy a vizsgálat egyértelművé teszi, az emberi tényező jelentős szerepet játszik az online biztonságban. Míg a felhasználók alapvetően tudatosak bizonyos szempontok alapján, például az információmegosztás kontrollálásában, a proaktív tájékozódás és a technológiai eszközök használata terén jelentős hiányosságok mutatkoznak. A bevezetésben felvetett problémák így az eredményekben is tükröződnek, alátámasztva, hogy az oktatás, a tájékoztatás és a technológiai eszközök alkalmazása együttesen szükséges a hatékony védelemhez.

Az eredmények alapján fontos az oktatás és a tudatos viselkedés fejlesztése, különös tekintettel a következőkre:

- Tájékoztatottság növelése: az adatvédelmi irányelvek és a biztonsági frissítések tartalmának megismerése.

- ♦ Technológiai eszközök használata: jelszókezelők és többfaktoros hitelesítés szélesebb körű alkalmazása.
- ♦ Közösségi szerepvállalás: kiberbiztonsági tudatosságot fejlesztő programok és kampányok támogatása.

Az eredmények megerősítik, hogy az egyéni viselkedés formálása és a technológiai megoldások együttes alkalmazása a leghatékonyabb módja az online visszaélések elleni védekezésnek.

3. AZ AKTÍV VÉDEKEZÉS ÖSSZEFÜGGÉSEI A DEMOGRÁFIAI HÁTTERREL

A kiberbiztonság az információs társadalom egyik központi kérdésévé vált, hiszen a digitális fenyegetések – például adathalászat, zsarolóvírusok és rosszindulatú programok – jelentős hatást gyakorolnak a gazdaságra, a magánéletre és az online kommunikáció biztonságára. A technológia fejlődése és az internet globális elterjedése új típusú sebezhetőségeket teremtett, amelyek kezelése proaktív megközelítést igényel mind az egyének, mind a szervezetek részéről.²⁴ Azonban nem minden felhasználó reagál egyformán ezekre a kihívásokra, mivel a demográfiai tényezők, például az életkor, nem, iskolai végzettség, gazdasági helyzet és földrajzi elhelyezkedés jelentős mértékben befolyásolják az egyének kiberbiztonsági tudatosságát és viselkedését.²⁵ A demográfiai tényezők figyelembevétele segít abban, hogy az egyes csoportokat megfelelő eszközökkel lássuk el, és növeljük az egyének tudatosságát és védekezési képességét.

A demográfiai tényezők szerepe a kiberbiztonságban azért is kulcsfontosságú, mert az egyes csoportok kockázati kitettsége és biztonságtudatossága jelentősen eltérhet. A következő szakaszokban részletesebben bemutatjuk azokat a demográfiai csoportokat, amelyeknél jellemzőbb az aktív védekező viselkedés, valamint azokat, akik különösen veszélyeztetettek a kiberbiztonsági fenyegetésekkel szemben. Az aktív védekezés a definíciója szerint olyan magatartást jelent, amely proaktívan reagál a potenciális támadásokra, beleértve a biztonsági rendszerek használatát, az oktatást és a tudatos viselkedési mintákat.

Az életkor meghatározó szerepet játszik a kiberbiztonsági attitűdök alakulásában. A fiatalabb generációk, akik „digitális bennszülöttként” nőnek fel, hajlamosak nagyobb arányban alkalmazni fejlett védelmi megoldásokat, miközben a közösségi média intenzív használatával sebezhetőséget is

²⁴ WOLFF 2021.

²⁵ DODEL–MESCH 2019.

teremtenek.²⁶ Ezzel szemben az idősebb korosztály alacsonyabb szintű technológiai jártassága és biztonsági tudatossága gyakran vezet a fenyegetésekkel szembeni kiszolgáltatottsághoz.²⁷

A nemi különbségek szintén jelentős tényezőt jelentenek: A San Diegó-i Kaliforniai Egyetem kutatói szerint a férfiak hajlamosabbak kockázatvállaló magatartásra az online térben, míg a nők gyakrabban összpontosítanak a személyes adatok és családtagjaik védelmére.²⁸ A gazdasági helyzet pedig nemcsak az eszközhasználat színvonalát, hanem a hozzáférhető kiberbiztonsági eszközök körét is meghatározza, így az alacsonyabb jövedelemmel rendelkező csoportok különösen kiszolgáltatottá válhatnak.²⁹

Ezen demográfiai tényezők vizsgálata nemcsak az egyéni kockázati szintek megértését teszi lehetővé, hanem segít abban is, hogy célzott oktatási programokat és védekezési stratégiákat dolgozzunk ki a különböző demográfiai csoportok számára. Az aktív védekezés, mint például a kétfaktoros hitelesítés, a jelszókezelők használata és a rendszeres frissítések, fontos eszközt jelentenek a kiberbiztonsági fenyegetésekkel szembeni ellenállás növelésében. A tanulmány célja annak ismertetése, hogy hogyan befolyásolják a demográfiai tényezők a kiberbiztonsági viselkedést, és bemutassa, hogy mely csoportok a leginkább veszélyeztetettek.

AZ AKTÍV VÉDEKEZŐ VISELKEDÉS JELLEMZŐI ÉS ELTERJEDÉSE KÜLÖNBÖZŐ DEMOGRÁFIAI CSOPORTOKBAN

Fiatalabb korosztály (18–35 év)

A 18–35 éves korosztály, amely magában foglalja az úgynevezett „digitális bennszülötteket,” jellemzően erősebb technológiai ismeretekkel rendelkezik, és jobban tisztában van a kiberfenyegetésekkel, mint az idősebb generációk.

²⁶ KIRSCHNER 2024.

²⁷ MARIANO et al. 2021.

²⁸ HARRIS–JENKINS 2006.

²⁹ BAUER 2018.

Ez a generáció gyakran napi szinten használ különféle digitális platformokat és eszközöket, beleértve a közösségi médiát, az online vásárlást és a *home office* során alkalmazandó megoldásokat. Ezek az aktivitások nagyobb kockázatot hordoznak magukban, és fokozott figyelmet igényelnek a biztonsági intézkedések betartására.

A fiatalabb felhasználók körében általában gyakoribb az olyan aktív védekező módszerek alkalmazása, mint a kétfaktoros hitelesítés, a jelszókezelő programok, a VPN-ek használata, valamint a rendszeres szoftverfrissítések. Egyes kutatások szerint a 18–35 éves korosztály tagjai nagyobb hajlandóságot mutatnak az új technológiai megoldások és biztonsági eszközök kipróbálására, valamint gyakrabban keresnek információkat a legújabb kiberbiztonsági fejlesztésekről. Ezzel szemben a fiatalok között a kockázatos magatartás is gyakori: a túlzott megosztás a közösségi médián vagy a kevésbé biztonságos jelszó szintén gyakori, ami növeli a támadási kitettséget.³⁰

Magasabb iskolai végzettséggel rendelkezők

Az iskolai végzettség szintén jelentős tényező a kiberbiztonsággal kapcsolatos tudatosság és a védekezési hajlandóság tekintetében. Azok, akik felsőfokú végzettséggel rendelkeznek, különösen, ha technológiai, mérnöki vagy informatikai területen dolgoznak, általában nagyobb kiberbiztonsági ismeretekkel rendelkeznek. Ez a csoport jobban érti a kiberfenyegetések természetét, és gyakrabban alkalmaz fejlett védelmi megoldásokat. Az oktatási szint emelkedésével növekszik az adathalász támadások és a közösségi média profilhamisítás elleni védekezés iránti érdeklődés is.

Az oktatott felhasználók nagyobb arányban ismerik és használják a biztonsági beállításokat, például a tűzfalakat, a multifaktoros hitelesítést és a felhőalapú biztonsági megoldásokat. Ezek az intézkedések segítik őket abban, hogy proaktívan reagáljanak a fenyegetésekre, és minimalizálják a támadások kockázatát.³¹

³⁰ VANNUCCI et al. 2020.

³¹ AL ABDULWAHID et al. 2015.

Férfiak és technológiai szektorban dolgozók

A férfiak körében, különösen azok között, akik az informatikai vagy technológiai szektorban dolgoznak, magasabb az aktív védekező viselkedés aránya. A férfiak technológiai érdeklődése és a technológiai ismeretek iránti nagyobb nyitottsága elősegíti a kiberbiztonsági tudatosságot, és az aktív védekezési módszerek széles körű alkalmazását. A technológiai iparágban dolgozók rendszerint mélyebb ismeretekkel rendelkeznek a kiberbiztonság terén, így jobban felkészültek a lehetséges támadások megelőzésére és kezelésére.³²

VESZÉLYEZTETETTEBB DEMOGRÁFIAI
CSOPORTOK A KIBERBIZTONSÁG TERÜLETÉN

Idősebb korosztály (55 év felett)

Az idősebb generációk általában kevésbé jártasak a digitális technológiák világában, ami sérülékenyebbé teszi őket a kibertámadásokkal szemben. E korcsoport tagjai gyakran kevésbé tájékozottak a kiberbiztonsági alapelvek és a digitális önvédelem eszközeinek terén. Számukra különösen nagy veszélyt jelentenek olyan támadások, mint az adathalász kampányok, a rosszindulatú kódot tartalmazó e-mailek és a zsarolóvírusok.

A digitális kockázatok észlelése az idősebbeknél alacsonyabb szintű lehet, és sokan közülük nem szívesen tanulnak új biztonsági technikákat vagy kezelik a szoftverfrissítéseket, amik pedig a védekezés alapvető részei. Kutatások szerint az idősebbek hajlamosak kevésbé biztonságos jelszavakat használni, és ritkábban alkalmaznak kétfaktoros hitelesítést, így könnyebben válhatnak a támadások célpontjává.³³

³² LI et al. 2019.

³³ ARON 2012.

Alacsonyabb gazdasági státuszú felhasználók

Az alacsonyabb jövedelemmel rendelkező felhasználók, akiknek korlátozottabbak az erőforrásai, gyakran kevésbé képesek megfelelő védelmi eszközöket igénybe venni. Számukra az olyan szolgáltatások, mint a fizetős vírusirtó szoftverek, a VPN-ek vagy a rendszeres adatmentések gyakran költségesek és nehezen elérhetők. A megfelelő védelmi eszközök hiánya miatt ők kiszolgáltatottabbak a támadásokkal szemben, és nagyobb az esélye annak, hogy személyes adataik veszélybe kerülnek.³⁴

Nők

Kiberbiztonsági szempontból a nemek közötti különbségek is megfigyelhetők. A nők általában hajlamosabbak nagyobb figyelmet fordítani a közösségi média használatára és a személyes adatok védelmére. Ugyanakkor a technológiai ismeretek hiánya miatt sok nő kevésbé hajlandó bonyolultabb biztonsági beállításokat alkalmazni, ami növeli a támadások kockázatát. Különösen kiszolgáltatottak a közösségi média hamis profiljainak és az online zaklatásoknak.³⁵

CÉLZOTT KIBERBIZTONSÁGI MEGOLDÁSOK A VESZÉLYEZTETETT DEMOGRÁFIAI CSOPORTOK SZÁMÁRA

A demográfiai tényezők figyelembevétele a kiberbiztonsági stratégiákban segít abban, hogy a felhasználók saját igényeikre és lehetőségeikre szabott védekezési módszereket alkalmazzanak. Az életkor, iskolai végzettség, nem és gazdasági státusz mind befolyásolja a felhasználók kockázati kitettségét és a védekezési hajlandóságot. A célzott kiberbiztonsági oktatás és az egyszerűen hozzáférhető védelmi eszközök segíthetnek abban, hogy a különböző

³⁴ KHAN–IKRAM–SALEEM 2024.

³⁵ MARWICK–MILLER 2014.

demográfiai csoportok is hatékonyabban védekezhessenek a digitális fenyegetésekkel szemben.

A kiberbiztonság szempontjából a legveszélyeztetettebb csoportok azok, akik demográfiai sajátosságaikból fakadóan fokozott kockázatnak vannak kitéve, és gyakran korlátozott hozzáféréssel rendelkeznek a szükséges védelmi eszközökhöz és ismeretekhez. Ezek a csoportok a társadalom különböző rétegeiből kerülnek ki, és sajátos kihívásokkal szembesülnek a kiberbiztonság terén, ami még nagyobb figyelmet igényel a döntéshozók és a szakértők részéről. Az alábbiakban bemutatjuk a legveszélyeztetettebb demográfiai csoportokat, az őket érintő specifikus kockázatokat és azokat az akadályokat, amelyek megnehezítik számukra a megfelelő védelem elérését.

Az eltérő demográfiai csoportok kiberbiztonsági védelme érdekében fontos, hogy testre szabott oktatási programokat és hozzáférhető védelmi eszközöket dolgozzanak ki a szakemberek. Az idősebb generációk számára például kiemelten fontos az egyszerűbb felhasználói felületek biztosítása, amelyek megkönnyítik a biztonsági funkciók elérését. Az alacsonyabb jövedelműek számára lényeges, hogy költséghatékony megoldások álljanak rendelkezésre, például ingyenes vagy kedvezményes áron elérhető alapvető vírusirtó szoftverek.

A fiatalabbak számára, akik hajlamosabbak a kockázatos viselkedésre, célszerű lenne interaktív oktatási anyagokkal felhívni a figyelmet a kiberbiztonsági fenyegetésekre. A technológiai háttérűeknek pedig szükségük lehet továbbképzésekre, amelyek a legújabb támadási formákkal és védelmi stratégiákkal ismertetik meg őket.

AZ EGYES DEMOGRÁFIAI CSOPORTOK VISELKEDÉSI JELLEMZŐI

Az egyes demográfiai csoportok eltérő viselkedési mintákat mutatnak a kiberbiztonság területén, ami a technológiai tapasztalatukból, életmódjukból, motivációikból és hozzáférési lehetőségeikből ered. Az alábbiakban bemutatjuk a különböző csoportok legjellemzőbb viselkedéseit, különös

figyelmet fordítva a kiberbiztonsági attitűdökre, védelmi intézkedésekre és a kockázati szokásokra.

FIATALABB KOROSZTÁLY (18–35 ÉV)

A fiatalabb korosztály, amely jellemzően „digitális bennszülöttként” szocializálódott, napi szinten használja az internetet és számos digitális eszközt. Az online jelenlét intenzitása és a technológia iránti nagyfokú érdeklődés jellemzi őket, ami különböző kiberbiztonsági magatartási formákban nyilvánul meg.

Jellemző viselkedések

- Aktív védekezési módszerek alkalmazása: A fiatalok körében gyakori az erős jelszavak használata, valamint a kétfaktoros hitelesítés és jelszókezelők alkalmazása. Mivel ők tudatosabbak a kiberbiztonsági kockázatokkal kapcsolatban, hajlamosabbak követni a biztonsági ajánlásokat.
- Kockázatos közösségimédia-magatartás: Habár tisztában vannak az alapvető biztonsági intézkedésekkel, a fiatalok hajlamosak túl sok személyes információt megosztani a közösségi médiában. A túlzott megosztás – például tartózkodási hely, napi rutin, kapcsolatok – fokozhatja a kiberbiztonsági kitettséget.
- Rendszeres eszköz- és szoftverfrissítés: A fiatalabb generáció tisztában van a szoftverek és operációs rendszerek frissítésének fontosságával, így gyakrabban frissítik eszközeiket, mint más korcsoportok.
- VPN-ek és adatvédelem iránti érdeklődés. A fiatalok körében egyre nagyobb figyelem irányul az adatvédelemre és a privát internet-használatra, ezért sokan használnak VPN-eket és más adatvédelmi megoldásokat.

IDŐSEBB KOROSZTÁLY (55 ÉV FELETT)

Az idősebb generációk, különösen az 55 év felettiek, a kiberbiztonság szempontjából kiemelten veszélyeztetett csoportnak számítanak. Ennek okai a technológiai ismeretek hiánya, a változásokkal szembeni esetleges ellenállás, valamint az új technológiák kezelésével kapcsolatos nehézségek. Ezt a korosztályt különösen gyakran célozzák meg olyan támadásokkal, mint az adathalászat, a telefonos csalások és a rosszindulatú e-mailek, amelyek az idősebb emberek nagyobb bizalmát és kisebb technológiai jártasságát használják ki.

Kiberbiztonsági kihívások az idősebbek körében

- ♦ Technológiai tudás hiánya: Sok idősebb ember nem ismeri a digitális önvédelem alapvető eszközeit, például a kétfaktoros hitelesítést, a jelszókezelőket vagy a vírusirtó szoftvereket. Ez a technológiai ismerethiány sebezhetővé teszi őket az online támadásokkal szemben.
- ♦ Alacsonyabb bizalmi küszöb: Az idősebb generációk általában nagyobb bizalommal viszonyulnak az e-mailekhez és telefonos kommunikációhoz, kevésbé gyanakodnak az ismeretlen üzenetekre vagy hívásokra, ezért gyakrabban esnek adathalászati támadások áldozatául.
- ♦ Digitális írástudás fejlesztésének szükségessége: Az idősebb korosztály számára fontos lenne a digitális írástudás fejlesztése, különösen olyan témákban, mint az e-mailes csalások felismerése, az erős jelszavak használata, valamint az adathalász webhelyek azonosítása.

Javasolt kiberbiztonsági intézkedések
az idősebb korosztály számára

Az idősebbek számára a kiberbiztonsági oktatási programok szervezése alapvető fontosságú, különösen a közösségi központokban és a helyi könyvtárakban, ahol egyszerű és gyakorlatias oktatási anyagokat kaphatnak.

A nagyobb átláthatósággal és könnyen érthető felületekkel rendelkező biztonsági alkalmazások szintén hozzájárulhatnak a védelem javításához ebben a csoportban.

Jellemző viselkedések

- Kevésbé biztonságos és egyszerű jelszavak használata: Az idősebbek gyakran egyszerű, könnyen megjegyezhető jelszavakat választanak, és hajlamosak ugyanazt a jelszót több webhelyen is használni, ami növeli a támadások kockázatát.
- Óvatlan e-mail-kezelési szokások: Gyakran kevésbé gyanakvók az ismeretlen forrásból származó e-mailekkel és linkekkel szemben, ami sebezhetővé teszi őket az adathalász támadásokkal és a rosszindulatú programokkal szemben.
- Korlátozott eszközhasználat és elavult szoftverek: Az idősebb generáció gyakran használ régebbi, elavult eszközöket, és kevésbé tartja fontosnak a rendszeres frissítést, ami így biztonsági réseket hagyhat nyitva
- Alacsonyabb adatvédelmi tudatosság: Az idősebbek körében kevésbé jellemző a privát böngészési módok vagy a VPN használata, mivel gyakran nem ismerik fel ennek fontosságát.

ALACSONYABB GAZDASÁGI STÁTUSZÚ FELHASZNÁLÓK

Az alacsony jövedelmű csoportok számára a kiberbiztonsági eszközök és megoldások gyakran anyagi akadályok miatt nehezen elérhetők. Ez a csoport kevésbé tájékozott a digitális önvédelem lehetőségeiről, és általában korlátozott hozzáféréssel rendelkezik a modern digitális eszközökhöz és az internetes oktatási forrásokhoz. További kockázatot jelent, hogy számukra gyakran megfizethetetlenek a prémium kiberbiztonsági megoldások, mint például a fizetős vírusirtók, a VPN-szolgáltatások és a speciális

adatvédelmi eszközök. Emellett az alacsony jövedelmű háztartások gyakran használnak régebbi, sebezhetőbb eszközöket, amelyeken már nem futnak a legújabb biztonsági frissítések.

Kiberbiztonsági kihívások az alacsony jövedelmű csoportok körében

- ♦ Elavult eszközök és szoftverek: A régebbi eszközökön és szoftvereken nem állnak rendelkezésre a legfrissebb biztonsági frissítések, ami növeli a rosszindulatú támadások sikerének esélyét. A költségek miatt ezek az emberek ritkán frissítik az eszközeiket, vagy használnak professzionális védelmi megoldásokat.
- ♦ Kiberbiztonsági eszközök elérhetetlensége: Az alacsonyabb gazdasági státuszú felhasználók számára kevésbé hozzáférhetők a fizetős biztonsági megoldások, így gyakran ingyenes, de limitált védelmi szintet biztosító programokat használnak.
- ♦ Oktatás és tudatosság hiánya: Sok alacsony jövedelmű ember számára nem állnak rendelkezésre a kiberbiztonsági oktatási források, és nincs megfelelő tudatosságuk a kiberfenyegetések felismeréséhez.

Javasolt kiberbiztonsági intézkedések az alacsony jövedelmű csoportok számára

Az ingyenesen vagy kedvezményesen elérhető, hatékony kiberbiztonsági eszközök kifejlesztése és terjesztése alapvető fontosságú lenne. Ezenkívül hatásosak lennének a kormányzati és nonprofit szervezetek által finanszírozott kiberbiztonsági oktatási kampányok és programok.

Jellemző viselkedések

- Ingyenes, de korlátozott védelmi szintet biztosító programok használata: Anyagi helyzetük miatt sokan csak ingyenes vírusirtót vagy korlátozott védelmi szintet nyújtó alkalmazásokat használnak, ami kevésbé hatékony védelmet jelenthet.
- Ritka frissítések és elavult hardverhasználat: Az anyagi korlátok miatt gyakran régebbi eszközöket használnak, amelyek már nem kapják meg a legújabb biztonsági frissítéseket, így jobban ki vannak téve a biztonsági réseknek.
- Alacsonyabb szintű jelszóhasználat és adatvédelem: Az ilyen felhasználók jellemzően kevésbé biztonságos jelszavakat használnak, és nem fektetnek hangsúlyt az adatvédelemre, mivel kevésbé vannak tudatában a védekezés fontosságának.

GYERMEKEK ÉS TIZENÉVESEK

A fiatalabb korosztály különösen ki van téve az online veszélyeknek, mivel sokuk számára az internet a mindennapi szocializáció és szórakozás szerves része. Az online térben jelen lévő rosszindulatú szereplők gyakran kihasználják a gyermekek és tizenévesek tapasztalatlanságát, különösen a közösségi médián és az online játéklatformokon keresztül.

Kiberbiztonsági kihívások a gyermekek és tizenévesek körében

- Közösségi média és online játékok veszélyei: A gyermekek és tinédzserek gyakran találkoznak online zaklatással, adatlopással és más fenyegetésekkel a közösségi média és az online játékok során. Az ő figyelmük könnyen elterelhető, és kevésbé ismerik fel a potenciális veszélyeket.

- ♦ Ismerethiány és információszűrési képesség: A fiatalok gyakran nem rendelkeznek elegendő tudással ahhoz, hogy megkülönböztessék a megbízható és a gyanús online forrásokat, így könnyen adathalászat áldozataivá válhatnak.
- ♦ Online kapcsolatok kialakítása: A tizenévesek közösségi médián keresztül gyakran ismerkednek idegenekkel, ami növeli a személyes adatokkal való visszaélés kockázatát és a potenciális online zaklatás veszélyét.

Javasolt kiberbiztonsági intézkedések a gyermekek és tizenévesek számára

Az iskolákban, valamint a szülői irányítás mellett megvalósított kiberbiztonsági oktatás segíthet abban, hogy a fiatalok megtanulják az online térben rejlő veszélyek elkerülését. Számukra az internet főleg a szórakozás és a közösségi kapcsolatok ápolásának eszköze, aminek kapcsán kevésbé fordítanak figyelmet a biztonságtudatos viselkedésre, ez még jobban növeli a veszélyeztetettségüket.

Ezenkívül a közösségi média és az online játékl platformok biztonságos használatának szabályaira is nagyobb hangsúlyt kellene fektetni, például a személyes adatok megosztásának korlátozására.

Jellemző viselkedések

- ♦ Községi médián való túlzott megosztás: A gyermekek és tizenévesek gyakran osztanak meg túl sok információt a közösségi médián, például fotókat, helyzetmeghatározást és személyes adatokat, ami kiszolgáltatottá teszi őket.
- ♦ Ismeretlen személyekkel való kommunikáció: Gyakran alakítanak ki kapcsolatokat idegenekkel, ami potenciális veszélyt jelenthet, különösen az online zaklatás vagy a személyazonosság-lopás szempontjából.

- ♦ Gyenge jelszavas védelem: A gyermekek és fiatalok gyakran egyszerű jelszavakat használnak, és hajlamosak ugyanazt a jelszót alkalmazni több platformon is, ami sebezhetővé teszi őket.
- ♦ Online játékok kockázatai: Az online játékok során a gyermekek gyakran személyes adatokat osztanak meg vagy játékbeli kiegészítőket vásárolnak, ami lehetőséget adhat a csalóknak a pénzügyi adatok megszerzésére.

SZOCIÁLISAN ELSZIGETELT VAGY ELSZEGÉNYEDETT KÖZÖSSÉGEK

Az elmaradottabb területeken, különösen a vidéki és szociálisan elszigetelt közösségekben élő emberek gyakran kevésbé tudják elérni a kiberbiztonsági erőforrásokat és oktatást, ami szintén nagyban hozzájárul sebezhetőségükhöz. Az ilyen közösségekben gyakoriak az alacsony technológiai készségek, az elavult eszközök, és kevésbé hozzáférhetők az oktatási anyagok.

Kiberbiztonsági kihívások a szociálisan elszigetelt közösségekben

- ♦ Korlátozott internet-hozzáférés: Az elmaradottabb területeken élők számára nem mindig biztosított a gyors és megbízható internetkapcsolat, ami korlátozhatja az információhoz való hozzáférést és a frissítések letöltésének lehetőségét.
- ♦ Alacsony technológiai szint és ismeretek: Ezekben a közösségekben kevésbé gyakori a modern digitális eszközök használata, ami tovább növeli a kiberbiztonsági kitettséget.
- ♦ Kiberbiztonsági oktatás hiánya: Az elszigetelt közösségek gyakran kimaradnak a kiberbiztonsági oktatási programokból, így az itt élők nem tudnak tájékozódni a veszélyekről és a védekezési lehetőségekről.

Javasolt kiberbiztonsági intézkedések az elszigetelt közösségek számára

Ezek a közösségek jellemzően kevésbé tájékozottak a kiberbiztonságról, és gyakran nincs lehetőségük hozzáférni a kiberbiztonsági oktatáshoz vagy megfelelő eszközökhöz. Az elszigetelt közösségek számára alapvetően fontos lenne a digitális eszközökhöz és a kiberbiztonsági oktatáshoz való hozzáférés biztosítása. Az ingyenes vagy kedvezményes oktatási programok szervezése, az internet-hozzáférés bővítése, valamint a könnyen érthető biztonsági útmutatók terjesztése segíthet csökkenteni a kockázatokat ezekben a közösségekben.

A szociálisan elszigetelt vagy alacsony technológiai fejlettségi szinten álló közösségek körében szintén megfigyelhetők bizonyos kockázatos viselkedési minták.

Jellemző viselkedések

- ♦ Alacsony szintű digitális írástudás: Az ilyen közösségekben sokan nem ismerik a kiberbiztonsági alapelveket, és kevésbé tájékozottak az adathalászat, a személyazonosság-lopás vagy a rosszindulatú programok veszélyei kapcsán.
- ♦ Ritka frissítések és biztonsági beállítások hiánya: Gyakran használnak elavult szoftvereket.

A fenti csoportok mindegyike speciális kihívásokkal néz szembe a kiberbiztonság terén, ami indokolja a célzott oktatási programok és védelmi intézkedések kidolgozását. A legveszélyeztetettebbek, mint az idősebbek, az alacsonyabb gazdasági státuszúak, a fiatalabbak és az elszigetelt közösségek, különösen sebezhetőek, mivel nem rendelkeznek megfelelő

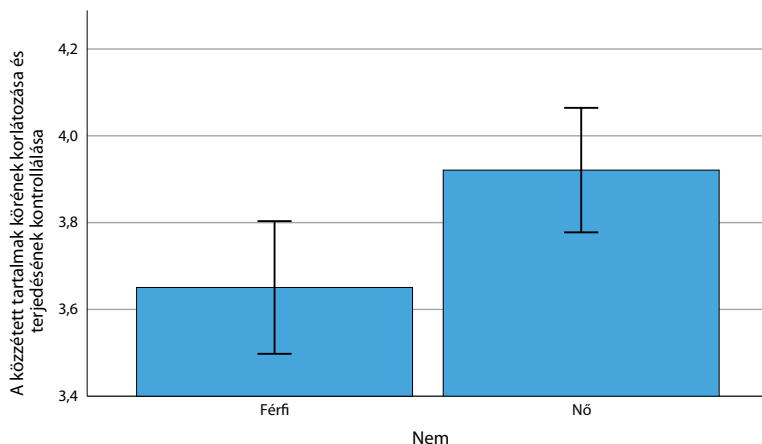
technológiai ismeretekkel vagy hozzáféréssel a védelmi eszközökhöz. Számukra az alapvető tudatosság növelése, valamint a hozzáférhető, költséghatékony biztonsági megoldások bevezetése elengedhetetlen ahhoz, hogy csökkenthessék a kiberbiztonsági fenyegetéseket, és biztonságosabban használhassák a digitális világot.

EREDMÉNYEK

Az aktív védekező viselkedések elterjedtségét feltáró jelleggel, valamennyi általunk mért demográfiai változó mentén elemeztük. E szakaszban ezek közül kizárólag azokat az elemzéseket mutatjuk be részletesen, amelyek esetében az elemzések szignifikáns különbségeket tártak fel különböző demográfiai csoportok között. A szignifikáns eredményeket nem mutató próbák részeredményeit nem közöljük, ezeket pusztán összefoglaló jelleggel említjük e szakasz végén.

Nemek közötti különbségek az aktív védekezésben

Elemzéseink azt mutatják, hogy a férfiak és nők között szignifikáns eltérések mutatkoznak az aktív védekező viselkedések gyakoriságában, ezek köre azonban a modellünk A_2 aldimenziójára korlátozódik, vagyis azokra a magatartási formákra, amelyek célja az online felületeken közzétett tartalmak korlátozása és azok terjedésének ellenőrzése ($t_{[213]} = -2,511, p = 0,013$). A 3.1. ábrán látható, hogy a két esetben a 95%-os konfidenciaintervallumok alig állnak fedésben egymással: a nőkre jellemzőbb, hogy jobban odafigyelnek arra, milyen információ kerül fel róluk a digitális térbe, mint a férfiakra, és többet is tesznek azért, hogy megelőzzék a privát információk eljutását illetéktelen személyekhez, illetve gazdasági társaságokhoz.



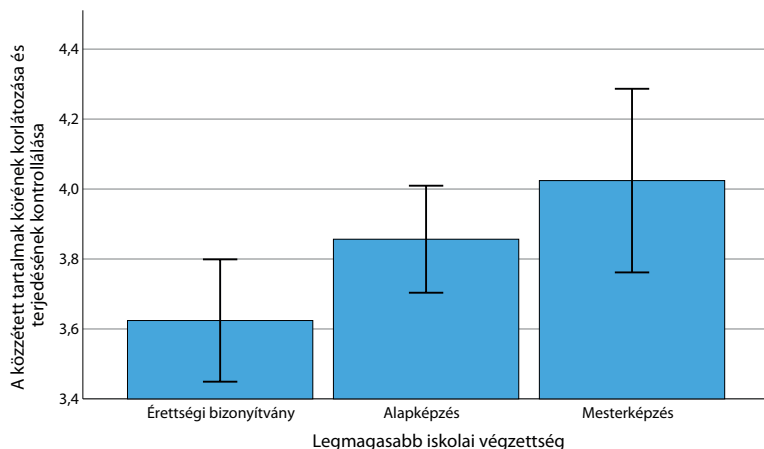
3.1. ábra: A nők szignifikáns előnye az online közzétett tartalmakhoz való hozzáférés ellenőrzésében (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

Bár a két nem között az eltérés mértéke csekély (Cohen-féle $d = -0,344$), a t -próba eredménye mégis azt mutatja, hogy e tendencia valós: a jelen mintanagyság mellett még ilyen kismértékű különbséget is csak igen ritka esetben okoz a mintavétel folyamatában szükségszerűen megjelenő véletlen.

Az aktív védekezés elterjedtsége különböző iskolai végzettségű csoportokban

Az elemzéseink (egyszempontos ANOVA-k) azt mutatják, hogy a védekező viselkedések gyakorisága az iskolai végzettséggel is szignifikánsan összefügg. E különbségek megmutatkoznak az aktív védekezés általános mutatójának szintjén ($F[2,208] = 4,014, p = 0,019$), a két fődimenzió közül azonban kizárólag a hozzáférés ellenőrzésével kapcsolatos viselkedésekben (A dimenzió; $F[2, 208] = 3,654, p = 0,028$), de még e dimenzióon belül is kizárólag



3.2. ábra: A különböző iskolai végzettségű csoportok közötti különbségek az online közzétett tartalmakhoz való hozzáférés ellenőrzésében (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

az A2 dimenzióban, vagyis ugyanazon viselkedések körében, amelyek esetében a nemek közötti eltéréseket is detektáltuk. E mintázat együttesen arról tanúskodik, hogy a végzettség alapvetően a közzétett tartalmak körének és terjedésének ellenőrzésével függ össze, és alapvetően e hatás mutatkozik meg a magasabb szintű viselkedéskategóriák szintjén. A 3.2. ábrán látható, hogy a magasabb iskolai végzettségű személyek általában többet tesznek azért, hogy az általuk közzétett tartalmak ne kerüljenek illetéktelen személyekhez vagy vállalatokhoz, és azt is alaposabban meggondolják, egyáltalán milyen információt tesznek közzé magukról online.

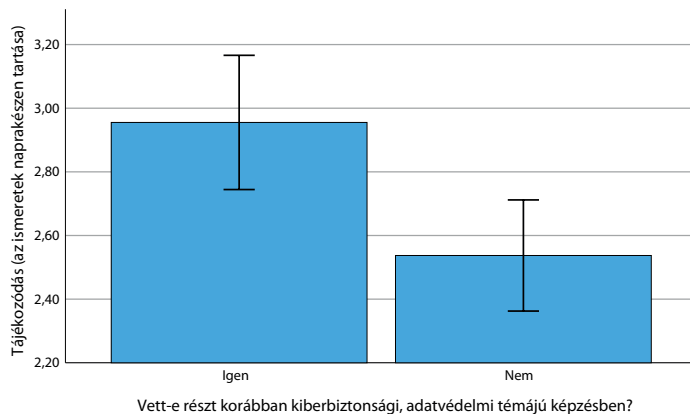
A szignifikáns eredményeket hozó varianciaanalíziseket lekövető *post hoc* elemzéseink eredményei szerint a különbség modellünk mindhárom szintjén (összesített mutató, A fődimenzió, A2 aldimenzió) kizárólag a két szélső csoport, vagyis a diplomával még nem rendelkezők és a mesterképzésen szerzett oklevéllel rendelkezők között szignifikáns. Az alapképzésen

szerzett oklevéllel rendelkező csoport konfidenciaintervalluma mindkét szélső csoporttal erősen átfed, ezért róluk biztosan nem állíthatjuk, hogy e viselkedések gyakorisága éppen esetükben a két szélső csoport átlagértéke közé esik. A leíró statisztika (mintaátlag) szintjén azonban az ábrán látható lépcsőzetes mintázat mindhárom mutató esetében egyaránt megfigyelhető.

A kiberbiztonsági képzések hatása a védekező viselkedésekre

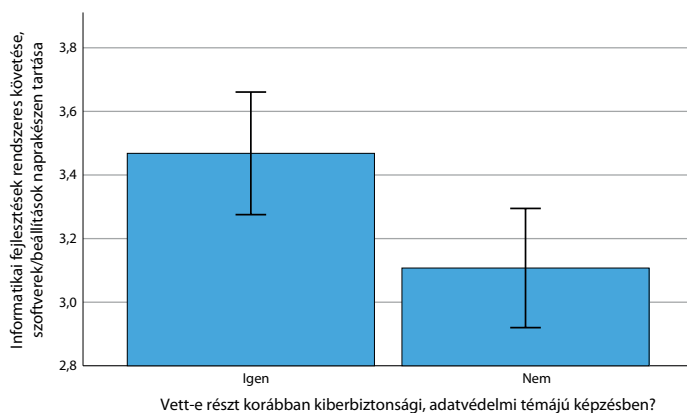
A harmadik demográfiai szempont, amely esetében szignifikáns hatást mutattak ki az elemzéseink, a kiberbiztonsági képzésben korábban részt vevő kitöltők és az ilyen jellegű tapasztalattal még nem rendelkező kitöltők összevetése volt. A két csoport közötti szignifikáns különbség megmutatkozik az általános, összesített mutató szintjén is ($t[213] = 2,639; p = 0,009$), de az specifikusabb próbák már azt mutatják, hogy e hatás kifejezetten a B fődimenzióban (naprakészség) detektálható különbségnek köszönhető ($t[213] = 3,278; p = 0,001$), míg az A fődimenzió esetében nem szignifikáns. Ráadásul a naprakészséghez kapcsolódó viselkedések körén belül az eltérés mindkét aldimenzió esetében egyaránt erősen szignifikáns (B1: $t[213] = 3,051, p = 0,003$; B2: $t[213] = 2,627, p = 0,009$). A 3.3. ábrán (B1 aldimenzió) látható, hogy a kiberbiztonsági képzésen részt vett személyek részletesebben tájékozódnak arról, hogyan változik az online környezet, milyen új fenyegetések jelennek meg, vannak-e hatékony védekezési módszerek e veszélyek ellen. A képzéseken átesett személyek ezen túlmenően részletesebben áttekintik az online felületeken jóváhagyott adatvédelmi irányelveket és a szerződések tartalmát is.

A B2 aldimenzióban megjelenő különbség pedig azt jelzi, e képzések hatására a résztvevők később jobban odafigyelnek arra, hogy eszközeiket is naprakészen tartsák, telepítsék a biztonsági frissítéseket, sőt, hogy rendszeresen ellenőrizzék: a kötelező frissítéseken túlmenően elérhető-e további biztonsági frissítések. E különbséget ábrázolja a 3.4. ábrán látható oszlopdiagram.



3.3. ábra: A kiberbiztonsági képzések szignifikáns hatása a rendszeres tájékozódásra (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés



3.4. ábra: A kiberbiztonsági képzések szignifikáns hatása az eszközök rendszeres frissítésére, naprakészen tartására (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

*Az aktív védekezéssel összefüggést nem
mutató demográfiai változók*

Az adatsorunkban szereplő további demográfiai változók egyike sem mutatott szignifikáns kapcsolatot az aktív védekező viselkedések egyik fő-, illetve aldimenziójával sem. Eredményeink szerint egyik viselkedéstípus gyakoriságának tekintetében sem mutathatók ki eltérések a különböző korcsoportok között: az életkor és a különböző szintű mutatók közötti korreláció minden esetben $r < 0,115$, amely a $p < 0,05$ szinten nem szignifikáns. Ehhez hasonlóan nem befolyásolja kimutatható mértékben az aktív viselkedések elterjedtségét, hogy az egyén milyen méretű településen él, aktív kereső-e, illetve felsőoktatási hallgató-e.

ÖSSZEGZÉS

Az eredmények rávilágítanak arra, hogy a demográfiai tényezők és az oktatás szignifikáns hatást gyakorolnak az egyének kiberbiztonsági viselkedésére és adatvédelmi tudatosságára. Az első eredmény szerint a nők nagyobb figyelmet fordítanak arra, hogy milyen információkat osztanak meg online, és többet tesznek adatvédelmük biztosításáért. Ez összhangban áll Harris és Jenkins kutatásával, amely kiemelte, hogy a nők általában óvatosabbak az online térben, mivel nagyobb hangsúlyt helyeznek a magánszféra védelmére.³⁶ Ugyanakkor ez az attitűd lehetőséget nyújt a nők számára célzott oktatási programok és egyszerűsített eszközök fejlesztésére, amelyek tovább erősíthetik a tudatosságukat a biztonság terén.

A második eredmény, miszerint a magasabb iskolai végzettségű személyek nagyobb odafigyeléssel kezelik az online megosztott tartalmaikat, szintén összhangban van a szakirodalommal. Dodel és Mesch megállapította, hogy az iskolázottabb egyének jobban tisztában vannak a digitális tér veszélyeivel, és elővigyázatosabban viselkednek.³⁷ Ez az eredmény alátámasztja azt a tényt, hogy az oktatás szintje nemcsak a kiberbiztonsági tudatosságot, hanem az adatokkal kapcsolatos felelős döntéshozatalt is erősíti. Az ilyen

³⁶ HARRIS–JENKINS 2006.

³⁷ DODEL–MESCH 2019.

felhasználók nagyobb arányban értékelik az adatvédelmi irányelvek átláthatóságát, és megfontoltabban kezelik az online interakcióikat.

A harmadik eredmény, amely szerint a kiberbiztonsági képzésen részt vett személyek részletesebb tudással rendelkeznek az új fenyegetésekről és védekezési stratégiákról, alátámasztja a célzott oktatás fontosságát. Li és szerzőtársai hangsúlyozták, hogy az oktatás és a képzés kulcsszerepet játszik a kiberbiztonsági kockázatok felismerésében és kezelésében.³⁸ A képzéseken részt vett személyek alaposabban vizsgálják az adatvédelmi irányelveket és szerződési feltételeket – ez egy olyan aktív védekezési forma, amely csökkentheti a támadások esélyét. Ez az eredmény arra is rámutat, hogy a rendszeres kiberbiztonsági képzés bevezetése szélesebb társadalmi rétegek számára hatékony eszköz lehet a digitális fenyegetésekkel szembeni ellenállás növelésére.

Összességében az eredmények megerősítik, hogy a nők adatvédelmi attitűdje, a magas iskolai végzettségűek körültekintő viselkedése és a képzések fontossága mind olyan tényezők, amelyek pozitív hatással vannak a kiberbiztonság általános szintjére. Ezek alapján a jövőbeli stratégiákban kiemelt figyelmet kell fordítani az adatvédelem és kiberbiztonsági képzés széles körű támogatására, különösen azok körében, akik kevesebb tapasztalattal vagy alacsonyabb iskolai végzettséggel rendelkeznek.

³⁸ LI et al. 2019.

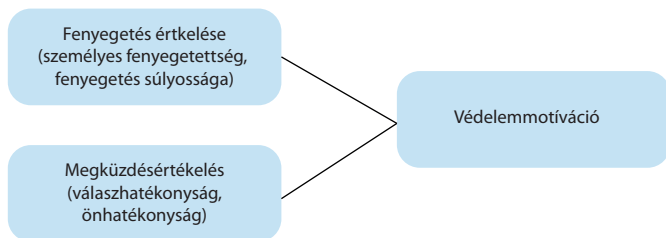
4. A VÉDELEMMOTIVÁCIÓ ÖSSZETEVŐI

A védelemmotiváció elméletét (*protection motivation theory*, a továbbiakban: PMT) eredetileg azért dolgozták ki, hogy megmagyarázza, hogyan reagálnak az emberek a félelmet keltő egészségügyi fenyegetésekkel kapcsolatos kommunikációra vagy „félelemfelhívásokra”. A védelmi motiváció arra a motivációra utal, hogy megvédjük magunkat egy egészségügyi fenyegetéssel szemben – ezt általában operatívan úgy határozzák meg, mint a javasolt cselekvés elfogadásának szándékát. A védelmi motivációt a fenyegetésértékelés és a megküzdési értékelési folyamat határozza meg: a fenyegetésértékelés olyan kognitív folyamat, amelyet az egyén a fenyegetettség szintjének (beleértve a súlyosságot, a sebezhetőséget és a félelmet) értékelésére használ, míg a megküzdési értékelési folyamat arra utal, hogyan értékeli az egyén a saját kockázatmegelőző attitűdjét, ami befolyásolja a védelmi motivációt (beleértve a válaszhatékonyágát és az önérvényesítést).³⁹ Így egy személy annál motiváltabb lesz a védekezésre (azaz erősebb a szándéka a javasolt cselekvés elfogadására), minél inkább úgy véli,

- hogy a fenyegetés valószínűsíthető, ha a jelenlegi cselekvést folytatja;
- hogy a fenyegetés bekövetkezése esetén a következmények súlyosak lesznek;
- hogy a javasolt cselekvés hatékonyan csökkenti a fenyegetés valószínűségét vagy súlyosságát;
- hogy képes végrehajtani a javasolt cselekvést.⁴⁰

³⁹ ALHAMAD–DONYAI 2021.

⁴⁰ SUTTON 2001.



4.1. ábra: A védelemmotiváció összetevői

Forrás: ALHAMAD–DONYAI 2021 alapján

A PMT tehát az emberi viselkedés pszichológiai megértésében kiemelkedő szerepet játszik, különösen a kockázatészlelés és a megelőző intézkedések kapcsolatának feltárásában. Az elméletet az egészségügy mellett az elmúlt évtizedben sikeresen alkalmazták az információbiztonság és kiberbiztonság területén is. A PMT alapját a fenyegetettségészlelés és a megküzdési mechanizmusok együttes értékelése képezi, amely meghatározza az egyének védelmi szándékait és viselkedését. Az elmélet különösen hasznosnak bizonyult a biztonságstudoatosság és a magatartás-változtatás elősegítésére irányuló kampányok tervezésében és végrehajtásában.⁴¹

Mivel a digitális fenyegetések és kiberbiztonsági problémák folyamatosan új kihívásokat jelentenek társadalmunk számára, a PMT segítségével könnyebb megértenünk, hogy az egyének hogyan reagálnak ezen fenyegetésekre, és mi motiválja őket a védekezési stratégiák alkalmazására. Az elmélet szerint a fenyegetettség észlelésének mértéke, a válaszadási lehetőségek hatékonyságába vetett hit, valamint az egyén saját képességei mind befolyásolják, milyen mértékben hajlandó az egyén védekezni.⁴² Az elmélet segít megérteni, hogy a fenyegetettség és a megküzdési képesség hogyan alakítja az egyének védelmi stratégiáit, különösen a digitális fenyegetésekkel szemben, mint például az adathalászat vagy a *malware* támadások.

A PMT egyik központi eleme a fenyegetésértékelés, amely az egyén által érzékelt fenyegetettség súlyosságából és a fenyegetésbe való beleesés

⁴¹ MOU et al. 2022.

⁴² HANUS–WU 2016.

valószínűségéből áll. A másik kulcsfontosságú tényező a megküzdési értékelés, amely az egyén válaszhatékonyágát és önhatékonyágát értékeli, így, ha az egyén úgy érzi, hogy képes hatékonyan reagálni a fenyegetésre, és hogy az alkalmazott védekezési stratégia hatékony, akkor nagyobb valószínűséggel alkalmazza az ajánlott viselkedést.⁴³

A PMT az egészségpszichológiai kutatásokban is jelentős szerepet kapott, különösen a Covid-19-pandémia alatt, ahol bebizonyosodott, hogy a PMT alkalmazása segíthet megérteni a közegészségügyi intézkedésekhez való alkalmazkodást, például a szociális távolságtartás és a maszkviselés elfogadását. Mindezek által a fenyegetés súlyosságának észlelése, valamint a válaszhatékonyág és önhatékonyág érzése előre jelezheti az egyének hajlandóságát az egészségügyi intézkedések betartására.⁴⁴

Az elmélet alkalmazása nemcsak az egészségügyi területeken, hanem a környezeti és társadalmi fenyegetésekre adott válaszok terén is sikeres. A PMT-t alkalmazták a környezeti kockázatokkal kapcsolatos magatartás előrejelzésére is, mint például a klímaváltozás hatásaira adott válaszok esetén. Az *augmented protection motivation theory* kiterjeszti az eredeti elméletet az irracionális viselkedés elemzésére is, például arra, hogy miért nem hajlandók az emberek a fenyegetettség ellenére védekezni, ha a kockázatokat észlelik.⁴⁵

A PMT-t az információs rendszerek és kiberbiztonság területén is alkalmazták. A kutatások szerint a kiberbiztonsági magatartás előrejelzésében a fenyegetettség és a megküzdési képességek kulcsfontosságúak, mivel az egyének védekezési szándékai gyakran függenek a digitális fenyegetések észlelésének mértékétől és a válaszingtezkedések hatékonyságától.⁴⁶

A kiberbiztonság területén a PMT hozzájárulása tehát különösen értékes, mivel segít megérteni, hogyan reagálnak az egyének a kiberfenyegetésekre, és milyen tényezők befolyásolják a védelmi viselkedésüket. Kutatások bizonyítják, hogy az egyének hajlandósága a biztonsági intézkedések betartására erősen függ a fenyegetés súlyosságának észlelésétől, a személyes

⁴³ CISMARU et al. 2011.

⁴⁴ YU-LAU-LAU 2022.

⁴⁵ MARIKYAN-PAPAGIANNIDIS 2023.

⁴⁶ HANUS-WU 2016.

sérülékenységük értékelésétől, valamint a javasolt védekezési stratégiák hatékonyságába vetett hitüktől. Ezenfelül a PMT alkalmazásával növelhetik a célzott tudatosító kampányok és képzések a dolgozók biztonságtudatos-ságát és felelősségtudatát az információbiztonság területén.⁴⁷

A PMT integrálása a kiberbiztonsági stratégiákba összességében nem csupán az egyéni viselkedés megértését segíti elő, hanem alapvetően hozzájárul a szervezeti szintű biztonsági kultúra fejlesztéséhez is. Az elmélet alkalmazásának tudományos és gyakorlati jelentősége igazolja, hogy a biztonsági tudatosság növelése és a fenyegetésekkel szembeni védekezés javítása elengedhetetlen eleme a megfelelő – egyaránt egyéni és társadalmi – kiberezhiliencia eléréséhez.⁴⁸

MÓDSZER: A VÉDELEMMOTIVÁCIÓ OPERACIONALIZÁCIÓJA

Ahhoz, hogy a védelemmotiváció érvényességét a kiberbiztonsági területen ellenőrizni tudjuk, ki kellett fejlesztenünk egy mérőeszközt, amely az elmélet minden konstrukcióját mérhetővé teszi, mégpedig oly módon, hogy a mérőeszköz tételei együttesen lefedik az online térben a felhasználóra leselkedő tipikus veszélyeket és az azok megelőzésére vagy a következmények enyhítésére leginkább alkalmas védekezési lehetőségeket.

A fenyegetés értékeléséhez tartozó két alkonstrukció esetében a tartalmi érvényesség biztosításához szükségünk volt egy olyan taxonómiára, amely sorra veszi és rendszerezi az online visszaélések jellemző típusait. Éppen ilyen struktúrában tárgyalja a kiberbűncselekmények sokrétű jelenségkörét Nurse tanulmánya,⁴⁹ amelyben a szerző a kiberbűnözés fajtáit öt nagy kategóriába sorolja: (1) szélhámosság és pszichológiai manipuláció (ez utóbira a magyar nyelvű szakirodalom is gyakran az angol *social engineering* kifejezéssel utal); (2) online zaklatás; (3) személyazonossággal kapcsolatos bűncselekmények; (4) hackelés; (5) szolgáltatástól,

⁴⁷ MENARD–BOTT–CROSSLER 2017.

⁴⁸ LEBEK et al. 2014.

⁴⁹ NURSE 2019.

információtól való elzárás. E kategóriák mindegyikét pontosan definiálja, és gazdag példaanyaggal illusztrálja. E kategóriák sorát a kérdőívfejlesztés során további két fenyegetéstípussal egészítettük ki, amelyek (6) az álhírekkel való megtévesztés, illetve (7) a mesterséges intelligenciával való visszaélések. E bővítést azért láttuk indokoltnak, mert kutatásunk körét nem korlátoztuk szigorú értelemben a jogilag és bűncselekménynek minősülő megtévesztésekre, ám a közösségimédia-platformokon gyakran terjednek olyan – érdeksérelmet vagy közriadalmat nem okozó – téves információk, amelyek terjesztése ugyan nem minősül bűncselekménynek, mégis az online térben megjelenő megtévesztés körébe tartozik. A különféle MI-alapú eszközök az elmúlt években pedig széles tömegek számára tették elérhetővé igen meggyőző, megtévesztő tartalmak létrehozását.

A fenyegetés értékeléséhez tartozó itemek szerkesztésekor tételesen végighaladtunk Nurse taxonómiáján, példáin, és életszerű, konkrét esetek leírásán keresztül igyekeztünk megragadni a különböző visszaéléstípusok leggyakoribb megnyilvánulási formáit. Ezen elv mentén konstruáltuk meg az utolsó két kategóriába sorolható visszaéléseket megjelenítő tételeket is.

A veszélyeztetettség érzetének mértékét egy Likert-mátrix segítségével mértük. E kérdőívblokkot az alábbi instrukció vezette be:

Az alábbiakban felsoroljuk a digitális eszközökkel elkövetett visszaélések néhány jellemző típusát. Kérjük, mindegyik eseményre vonatkozóan jelölje, mennyire tartja valószínűnek, hogy a következő 12 hónapban Ön személyesen áldozatává válik egy ilyen típusú visszaélésnek!

Amennyiben a múltban e visszaéléstípusok legalább egyikének már áldozata volt, akkor azt becsülje meg, mennyire tartja valószínűnek, hogy az adott esemény újra előfordul Önnel.

- 1 – Szinte biztos, hogy ilyen velem nem fog előfordulni.
- 2 – Valószínűtlen, hogy ez velem megtörténjen.
- 3 – Nem tudom eldönteni; lehet, hogy megtörténik velem, de lehet, hogy nem.
- 4 – Valószínű, hogy ez velem is megtörténik majd.
- 5 – Szinte biztos, hogy ez velem is meg fog történni.

A fenyegetés súlyosságát szintén Likert-mátrixszal mértünk. Ez esetben az instrukció a következő volt:

Az emberek különböző mértékben találják aggasztónak az online visszaélések lehetőségét. Például van, aki egy legyintéssel elintézi, ha a közösségi médián sértő üzeneteket kap, de van olyan felhasználó is, akinek ez súlyos lelki megterhelést jelent. Kérjük, most olvassa el újra a fent már látott felsorolást, de most nyilatkozzon arról, hogy az Ön életében mennyire jelentene súlyos problémát, ha az adott esemény valóban bekövetkezne!

- 1 – Nem aggódnék különösebben. Előfordul az ilyesmi: kezelném a helyzetet, és túltenném magam rajta.
- 2 – Kicsit bosszankodnék vagy aggódnék a következmények miatt, de azért vannak az életben ennél súlyosabb gondok.
- 3 – Ez jelentős probléma lenne az életemben, ijesztő belegondolni.
- 4 – Ez súlyos probléma lenne számomra, rendkívül rémisztő még a gondolat is, hogy ilyen előfordulhat.
- 5 – Ha velem ilyen megtörténne, úgy érzem, az katasztrofális következményekkel járna. Ez lenne az életem egyik legsúlyosabb problémája.

Ahogy a második blokk instrukciójából látható, e két szakaszban a tételek sora azonos volt, tehát az adatközlők ugyanazokról az eshetőségekről nyilatkoztak, megítélve, mennyire tartják valószínűnek, hogy ők maguk is az adott típusú megtévesztés áldozatává válnak, és mennyire súlyos problémát jelentene számukra, ha ez bekövetkezne.

Az egyes tételek pontos tartalmát és visszaéléstípusok taxonómiáját a 4.1. táblázat foglalja össze (a taxonómia kategóriáit az adatközlők nem látták, ez biztosította, hogy az egyes eseményeket egymástól függetlenül ítéljék meg, és válaszaikat ne befolyásolja, milyen általános típusba tartozik egy-egy visszaélés):

4.1. táblázat: A veszélyeztetettség érzetét és fenyegetés
súlyosságát mérő kérdőív tételek

Kategória	Tétel
Szélhámosság, pszichológiai manipuláció	Egy csaló megtéveszt azzal, hogy valamilyen szolgáltatónak (banknak, közüzemnek, futárszolgálatnak) adja ki magát, és rávesz, hogy megadjak magamról valamilyen személyes adatot.
	Egy csaló megtéveszt azzal, hogy személyes ismerősömnnek vagy ismerősöm barátjának adja ki magát, és így vesz rá valamire.
Online zaklatás	Valakitől ismétlődően bántó, sértő tartalmú üzeneteket kapok, amelyek a küllememet, a tulajdonságaimat vagy a képességeimet/hozzáértésemet becsmérlik.
	Egy online közösségben, amelynek tagja vagyok, valaki provokatív megjegyzéseivel szándékosan feszültséget gerjeszt, és engem is támadni kezd (trollkodás).
	Egy ismeretlen személy érzékeny információkat kezd gyűjteni rólam, követi az online tevékenységemet, és mindezt a tudomásomra is hozza, hogy ezáltal félelmet keltsen bennem, és ne érezzem magam biztonságban.
	Egy ismeretlen személy rágalmozó üzeneteket küld rólam az ismerőseimnek.
	Valaki azzal próbál zsarolni, hogy intim felvételekkel rendelkezik rólam, és ha nem teljesítem a feltételeit, ezeket eljuttatja rokonaimnak, barátainak és kollégáimnak.
Személyazonossággal kapcsolatos bűncselekmények	Tanúja leszek annak, hogy egy online üzenetben valaki egy olyan társadalmi csoport ellen buzdít erőszakra, amelynek én is tagja vagyok (például származásom, szexuális orientációm, vallásom vagy politikai szimpátiám alapján).
	Valaki megszerzi a banki adataimat, és a nevemben vásárol, átutal, hitelkártyát vált ki vagy kölcsönt vesz fel.
	Egy haragosom nyilvánosan közzéteszi a személyes adataimat (lakcímem, telefonszámom, e-mail-címem stb.), hogy így álljon rajtam bosszút.

Kategória	Tétel
Hackelés	A számítógépemre vagy okostelefonomra vírus vagy más rosszindulatú szoftver kerül.
	A számítógépemre vagy okostelefonomra kémprogram kerül, amely segítségével illetéktelen személyek hozzáférhetnek az eszközön tárolt adatokhoz.
	A számítógépemre vagy okostelefonomra olyan kémprogram kerül, amely megfigyeli és illetéktelen személynek közvetíti a képernyőn látható információkat vagy a billentyűleütéseket.
	A számítógépemre vagy okostelefonomra olyan kémprogram kerül, amely a tudtom nélkül felvételeket készít rólam a saját eszközőm kamerájával.
	Nyilvános helyen valaki – személyesen vagy kamera használatával – elles egy jelszót, amelyet begépelek, így hozzáfér bizonyos adataimhoz.
Szolgáltatástól vagy információtól való elzárás	Valaki – a rólam online elérhető információk alapján – egyszerű próbálkozással, találgatással kitalálja az egyik jelszavam, így hozzáfér bizonyos adataimhoz.
	Egy zsarolóvírus titkosítja és hozzáférhetetlenné teszi a számítógépen tárolt adatokat, és az elkövető váltságdíjat követel a titkosítás feloldásáért.
Álhírek	Olyan hackertámadás ér egy általam is igénybe vett szolgáltatást (például bank, tanulmányi rendszer, ügyfélkapu), amely számomra is bosszúságot vagy kárt okoz.
	Egy hivatalos hírmédium által közzétett hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.
Mesterséges intelligenciával való visszaélések	Egy közösségi médiában terjedő, bizonytalan eredetű hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.
	Egy mesterséges intelligencia által generált, személyesen nekem címzett szöveges üzenetről elhiszem, hogy azt ember (például ügyfélszolgálati munkatárs) fogalmazta.
	Egy mesterséges intelligencia által generált fényképről elhiszem, hogy az valódi.

Forrás: saját szerkesztés

A válaszhatékonyásra vonatkozó hiedelmeket külön szakaszban mértük, amelyet az alábbi magyarázat és instrukció vezetett be:

Az emberek eltérően vélekednek arról, milyen hatékonyan lehet különféle módokon védekezni az online visszaélések ellen. A leghatékonyabb védelmet persze az jelenti, ha a felhasználó sokféle módon védekezik, de az egyes módszerek hatékonyságáról megoszlanak a vélemények.

Kérjük, jelölje, Ön mennyire tartja az alábbi módszereket hasznosnak és fontosnak az online visszaélések elleni védekezésben. Az egyes értékek jelentése a következő:

- 1 – Ez nem lényeges, semmiféle védelmet nem jelent.
- 2 – Esetenként segíthet, de a tapasztalt elkövetők ellen nem jelent védelmet.
- 3 – Hasznos, de önmagában nem elegendő.
- 4 – Nagyon hasznos: ha erre figyelünk, ezzel már önmagában sok visszaélés megelőzhető.
- 5 – Az egyik leghatékonyabb módszer: bizonyos típusú visszaélések ellen szinte tökéletes védelmet nyújt.

A tételek megkonstruálásakor azon viselkedésformák körének lefedésére törekedtünk, amelyek a szakirodalom, a kiberbiztonsági képzések tananyaga, és az online médiában megjelenő ismeretterjesztő cikkek szerint az előző blokkban felsorolt visszaélések ellen nyújthatnak hatékony vagy részleges védelmet. A felsorolt védekezésmódok a 4.2. táblázatban láthatók.

4.2. táblázat: A válaszhatékonyaságot mérő kérdőívtételek

Védekezésmód
Intézménytől érkező üzenetek esetében mindig ellenőrizzük a feladót!
Ha egy hivatkozás intézményi weboldalra irányít, mindig ellenőrizzük az oldal URL-címét!
Ha régi vagy távoli ismerős veszi fel velünk a kapcsolatot, igyekezzünk ellenőrizni, hogy a személy valóban az-e, akinek mondja magát!
Online közösségekben, közösségi médiában ne foglazzunk meg olyan álláspontot, amely miatt gyűlölködés célpontjává válhatunk!
Ha valaki sértőt vagy gyűlölködő tartalmú üzeneteket küld, szakítsuk meg vele az online kapcsolatot (töröljük, tiltsuk le)!
Ne osszunk meg magunkról nyilvánosan érzékeny, személyes információkat!

Védekezés mód

Privát közösségimédia-profilat használjunk, amelyen kizárólag a beleegyezésünkkel válhat valaki a követőnké, ismerősünké!

Takarjuk le a számítógépünk kameráját, ha az nincs használatban!

Videóhívás közben ne tegyünk olyat, amiről ránk nézve kínos vagy kétértelmű képernyőkép készíthető!

Senkinek ne engedjük, hogy olyan fénykép- vagy videófelvételeket készítsen rólunk, amelyekkel később bárki más visszaélhet!

Ne látogassunk olyan weboldalakat, amelyek ingyen kínálnak egyébként fizetős szolgáltatást!

A nem garantáltan megbízható weboldalak meglátogatásakor használjuk a böngészőnket inkognitó módban!

Rendszeresen töröljük a böngészőnk által eltárolt sütiket!

Ne töltsünk le jogvédett szoftvereket fájlmegosztó rendszerekről!

Ne töltsünk le egyéb jogvédett tartalmakat (például filmeket, könyveket) fájlmegosztó rendszerekről!

Ha nyilvános helyen kell egy fontos jelszót begépelnünk, győződjünk meg róla, hogy nincs senki a közelünkben, és takarjuk el a képernyőt/billentyűzetet, amennyire csak lehetséges!

Banki vagy más közületi alkalmazásból/honlapról az ügyintézés végén azonnal jelentkezünk ki!

Ne használjunk olyan jelszót, amely családtagjaink/házikedvencünk nevét, monogramját, születési évszámot, hobbinkkal kapcsolatos szót vagy más könnyen kitalálható információt tartalmaz!

Minden platformon más, egyedi jelszót használjunk!

A jelszavainkat rendszeresen változtassuk meg!

Ne használjunk olyan szolgáltatást, amely a jelszavaink elmentésével megkönnyíti a különféle fiókokba való belépést!

Személyes adatainkat (e-mail-cím, telefonszám) ne tegyük közzé nyilvános felületen!

Online vásárlásra tartsunk fenn külön számlát és bankkártyát, amelyen egyébként nem tartunk jelentős összeget!

Rendszeresen frissítsük a számítógépünk/okostelefonunk operációs rendszerét és a használt alkalmazásokat!

Ne bízzuk magunkat az operációs rendszer beépített vírusvédelmére: vegyünk igénybe professzionális vírusvédelmet adó szolgáltatást!

A számítógépünkön ne rejtjük el az állományok kiterjesztését (például „pdf”, „docx”, „docm”, „exe”)!

A fontos, határidős online ügyleteket (például banki átutalás, adóbevallás) ne halogassuk az utolsó pillanatra!

Az online olvasott hírek tartalmát mindig ellenőrizzük más forrásokban is!

A valószínűtlen helyzeteket/eseményeket ábrázoló fényképeken ellenőrizzük az apró részleteket, amelyek esetleg arról árulkodnak, hogy a felvétel nem valódi!

Forrás: saját szerkesztés

Végül a kérdőívblokk utolsó szakaszában az önhatékonyság konstrukcióját operacionalizáltuk. Az instrukció egyértelművé tette a kitöltő számára, hogy ezen a ponton nem az egyes védekezés módok hatékonyságáról általában, hanem saját képességeiről és önfegyelméről kell nyilatkoznia:

Sokan úgy érzik, kellő körültekintéssel képesek arra, hogy megvédjék magukat az online visszaélések ellen. Mások viszont kiszolgáltatottnak érzik magukat, mert úgy vélik, nem rendelkeznek az ehhez szükséges képességekkel vagy önfegyelmel.

Ön hogy érzi, mennyire képes az alábbi módokon védekezni az ellen, hogy áldozattá váljon?

- 1 – Egyáltalán nem.
- 2 – Kis mértékben.
- 3 – Változó.
- 4 – Többnyire.
- 5 – Teljes mértékben.

E szakaszban a kérdőív tételek olyan megküzdési stratégiákat írtak le, amelyekről feltételezhető, hogy – a hatékonyságukba vetett bizalomtól függetlenül – egyénenként változó lehet, ki milyen mértékben tartja alkalmasnak magát az adott védekezés mód fegyelmezett és szakszerű kivitelezésére. A konkrét tételek felsorolása a 4.3. táblázatban található.

4.3. táblázat: Az önhatékonyságot mérő kérdőív tételek (minden állítás a „Képes vagyok...” mondatkezdetet egészíti ki)

Megküzdési stratégia
...ellenőrizni, hogy egy üzenet hiteles forrásból származik-e.
...ellenőrizni, hogy egy intézmény valódi weboldalát látogattam meg, és nem annak hamisított másolatát.
...figyelni arra, hogy kizárólag olyan üzeneteket írjak online, amelyeket a későbbiekben is válllok, és nem érintene kínosan, ha nyilvánosságra kerülnének.
...figyelni arra, hogy ne osszak meg magamról online érzékeny, személyes információkat.

Megküzdési stratégia

- ...figyelni arra, hogy semmilyen körülmények között ne készülhesen rólam olyan fénykép- vagy videófelvétel, amelyről nem szeretném, ha nyilvánosságra kerülne.
 - ...arra, hogy a számomra szükséges jogvédett tartalmakat (például könyvek, filmek) mindig jogszerű forrásból szerezzem be, és ellenálljak az ingyenes, de nem jogszerű források kísértésének.
 - ...biztonságosan használni a böngészőprogramot.
 - ...a biztonságom érdekében átgondoltan tesztekre szabni a közösségimédia-profilom beállításait.
 - ...megkülönböztetni a biztonságos weboldalakat a gyanús és potenciálisan veszélyes honlapoktól.
 - ...a számítógépen megkülönböztetni a futtatható állományokat az adatokat tartalmazó állományoktól.
 - ...arra, hogy odafigyeljek az általam használt operációs rendszerek és egyéb szoftverek rendszeres frissítésére.
 - ...vírusirtó és egyéb biztonsági szoftverek telepítésére, használatára és rendszeres frissítésére.
 - ...vigyázni arra, hogy nyilvános helyen senki ne lássa az általam begépett jelszót/PIN-kódot.
 - ...arra, hogy a fontos felhasználói vagy ügyfélprofiljaimat biztonságos jelszóval védjem meg.
 - ...minden platformon eltérő jelszót használni.
 - ...a jelszavaimat biztonságosan kezelni.
 - ...garantálni az online vásárlásaim és banki adataim biztonságát.
 - ...arra, hogy ne halogassam a határidős, online eszközöket igénylő feladatokat az utolsó pillanatra.
 - ...arra, hogy megkülönböztessem az álhíreket a valódiaktól.
 - ...arra, hogy megkülönböztessem a mesterséges intelligencia által generált szövegeket az ember által fogalmazott szövegektől.
 - ...arra, hogy megkülönböztessem a mesterséges intelligencia által generált fényképeket a valódi fényképfelvételektől.
-

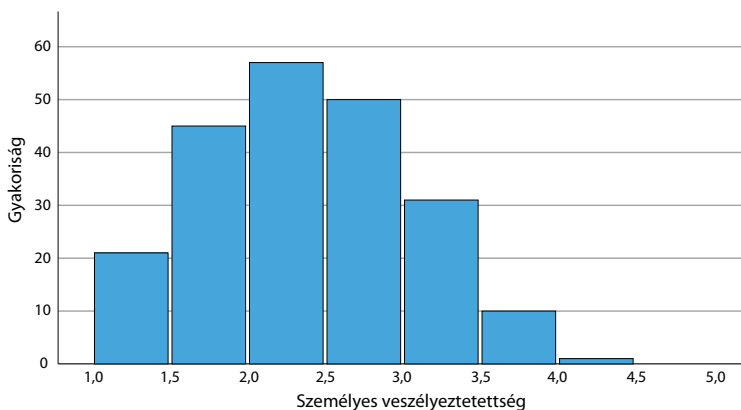
Forrás: saját szerkesztés

A védelemmotiváció mind a négy tényezője esetében az összesített mutató értékét a válaszok átlagértéke adta. A fenyegetés valószínűségének és súlyosságának mérésekor alacsonyabb szintű mutatókat is konstruáltunk, amelyek a hét visszaéléstípushoz tartozó itemekre adott válaszokat összesítik.

A későbbiekben az összefüggések feltárását célzó elemzésekben ezeket a specifikusabb mutatókat nem használtuk, de a jelen fejezetben számot adunk ezek eloszlásainak összevetéséről is.

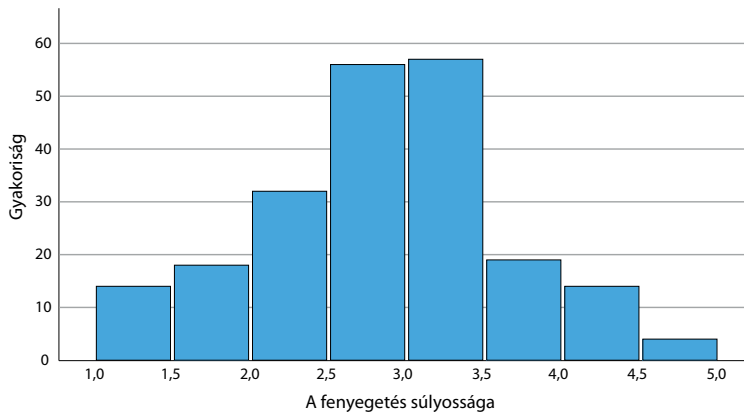
EREDMÉNYEK

A védelemmotiváció mind a négy összetevője esetében a méréseink megbízhatósági mutatói kiválóak: a Cronbach-féle alfa értéke a veszélyeztetettség, a súlyosság, a válaszhatékonyság és az önhatékonyság esetében rendre 0,92; 0,95; 0,94 és 0,92. Ahogy 4.2.–4.5. ábrákon szereplő hisztogramokon látható, mind a négy összesített mutató halom alakú és közelítőleg szimmetrikus eloszlást mutat. Bár a válaszhatékonyság és az önhatékonyság eloszlása kissé balra ferde, ezt a mintázatot viszonylag kis számú, nagymértékű negatív deviációjú érték okozza: ezektől eltekintve extrém ferdeség vagy lapultság egyik esetben sem figyelhető meg.



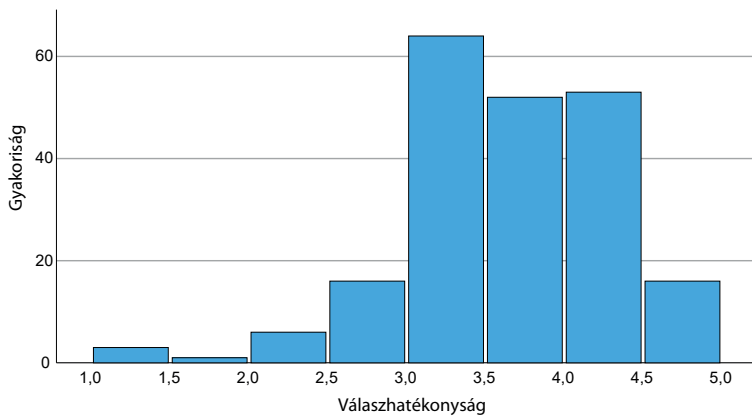
4.2. ábra: A személyes veszélyeztetettség fokát jelző összesített mutató eloszlása

Forrás: saját szerkesztés



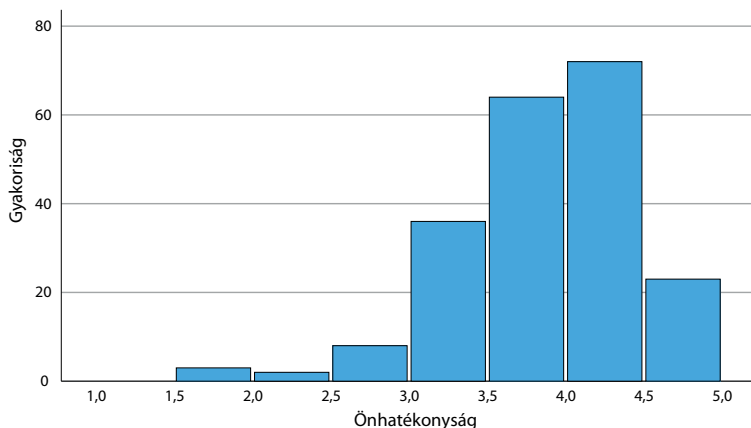
4.3. ábra: A fenyegetés súlyosságának érzetét jelző összesített mutató eloszlása

Forrás: saját szerkesztés



4.4. ábra: A válaszhatékonyság mértékét kifejező összesített mutató eloszlása

Forrás: saját szerkesztés



4.5. ábra: Az önhatékonyosság fokát kifejező összesített mutató eloszlása

Forrás: saját szerkesztés

A személyes veszélyeztetettség érzete

Bár a későbbi elemzéseinkben kizárólag az összesített mutatók játszanak majd szerepet, ezen a ponton érdemes a leíró statisztikák szintjén megvizsgálnunk az egyes tételek viselkedését is, hiszen ezek sok részletet árulnak el arról, milyen típusú visszaélésektől érzik magukat veszélyeztetve az adatközlők, melyek érintenék őket a legsúlyosabban és így tovább.

A személyes veszélyeztetettség esetében a tételek az átlagértékek szerint csökkenő rendben a 4.4. táblázatban láthatók.

*4.4. táblázat: Az online visszaélések valószínűségének
megítélése – a személyes veszélyeztetettség érzetét mérő tételek
leíró statisztikái az átlag szerint csökkenő rendben*

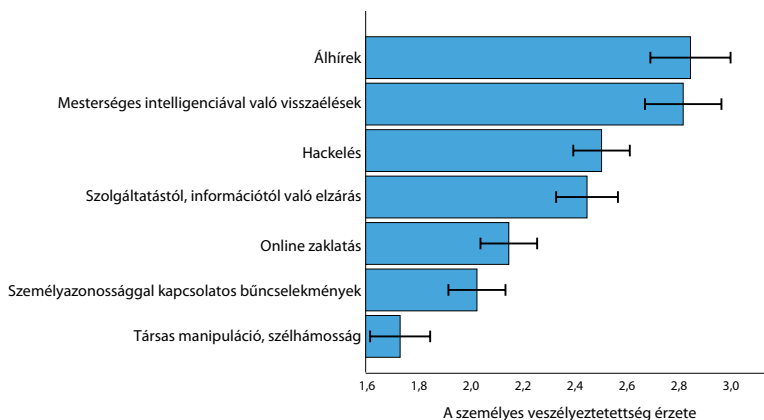
	Átlag	Szórás
Egy mesterséges intelligencia által generált fényképről elhiszem, hogy az valódi.	3,17	1,287
A számítógépre, esetleg okostelefonra vírus vagy más rosszindulatú szoftver kerül.	3,10	1,083
Egy hivatalos hírmédium által közzétett hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.	2,98	1,213
Olyan hackertámadás ér egy általam is igénybe vett szolgáltatást (például bank, tanulmányi rendszer, ügyfélkapu), amely számomra is bosszúságot vagy kárt okoz.	2,88	1,180
Egy közösségi médiában terjedő, bizonytalan eredetű hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.	2,72	1,252
A számítógépre vagy okostelefonra kémprogram kerül, amely segítségével illetéktelen személyek hozzáférhetnek az eszközön tárolt adatokhoz.	2,66	1,073
A számítógépre vagy okostelefonra olyan kémprogram kerül, amely megfigyeli és illetéktelen személynek közvetíti a képernyőn látható információkat vagy a billentyűleütéseket.	2,60	1,076
Egy online közösségben, amelynek tagja vagyok, valaki provokatív megjegyzéseivel szándékosan feszültséget gerjeszt, és engem is támadni kezd (trollkodás).	2,50	1,297
A számítógépre vagy okostelefonra olyan kémprogram kerül, amely a tudtom nélkül felvételeket készít rólam a saját eszközőm kamerájával.	2,48	1,063
Egy mesterséges intelligencia által generált, személyesen nekem címzett szöveges üzenetről elhiszem, hogy azt ember (például ügyfél-szolgálati munkatárs) fogalmazta.	2,47	1,195
Tanúja leszek annak, hogy egy online üzenetben valaki egy olyan társadalmi csoport ellen buzdít erőszakra, amelynek én is tagja vagyok (például származásom, szexuális orientáció, vallásom vagy politikai szimpátiám alapján).	2,39	1,383
Valaki megszerzi a banki adataimat, és a nevemben vásárol, átutal, hitelkártyát vált ki, vagy kölcsönt vesz fel.	2,28	1,062
Nyilvános helyen valaki – személyesen vagy kamera használatával – elles egy jelszót, amelyet begépelek, így hozzáfér bizonyos adataimhoz.	2,26	1,075

	Átlag	Szórás
Valakitől ismétlődően bántó, sértő tartalmú üzeneteket kapok, amelyek a küllememet, a tulajdonságaimat vagy a képességeimet/ hozzáértésemet becsmérlik.	2,21	1,188
Egy ismeretlen személy érzékeny információkat kezd gyűjteni rólam, követi az online tevékenységemet, és mindezt a tudomásomra is hozzá, hogy ezáltal félelmet keltsen bennem, és ne érezzem magam biztonságban.	2,20	1,027
Egy ismeretlen személy rágalmozó üzeneteket küld rólam az ismerőseimnek.	2,04	,951
Egy zsarolóvírus titkosítja és hozzáférhetetlenné teszi a számítógépen tárolt adatokat, és az elkövető váltságdíjat követel a titkosítás feloldásáért.	2,02	,952
Valaki – a rólam online elérhető információk alapján – egyszerű próbálkozással, találgatással kitalálja az egyik jelszavamot, így hozzáfér bizonyos adataimhoz.	1,94	1,040
Egy csaló megtéveszt azzal, hogy valamilyen szolgáltatónak (banknak, közüzemnek, futárszolgálatnak) adja ki magát, és rávesz, hogy megadjak magamról valamilyen személyes adatot.	1,82	1,134
Egy haragosom nyilvánosan közzéteszi a személyes adataimat (lakcímem, telefonszámom, e-mail-címem stb.), hogy így álljon rajtam bosszút.	1,78	0,857
Egy csaló megtéveszt azzal, hogy személyes ismerősömnek vagy ismerősöm barátjának adja ki magát, és így vesz rá valamire.	1,64	0,900
Valaki azzal próbál zsarolni, hogy intim felvételekkel rendelkezik rólam, és ha nem teljesítem a feltételeit, ezeket eljuttatja rokonaimnak, barátainak és kollégáimnak.	1,57	0,929

Forrás: saját szerkesztés

A rendezett lista felső és alsó szegmensét összevetve szembeötlő az a tendencia, hogy a kitöltők úgy gondolják, igen valószínű, hogy a közeljövőben mesterséges intelligenciával előállított tartalommal vagy álhírrrel sikeresen meg fogják őket tévesztetni, de nagy valószínűséggel bekövetkező eseménynek ítélik meg azt is, hogy hackertámadás áldozatai lesznek. Ezzel szemben nem tartják valószínűnek, hogy az online felületeken csalók vagy szélhámosok sikeresen rávegyék őket olyasmire, amit egyébként nem tennének meg.

Az egyes visszaéléstípusok szerint összesített mutatók sora is ezt a min-tázatot mutatja (4.6. ábra).



4.6. ábra: A veszélyeztetettség érzete visszaéléstípusok szerint összesítve
(a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

A fenyegetés súlyosságának megítélése

Amikor a kitöltőknek arról kellett nyilatkozniuk, mennyire érintené őket súlyosan, ha valóban áldozatává válnának a felsorolt visszaéléseknek, a mintázat csak részben tükrözte a valószínűség megítélésével kapcsolatos válaszaikat. A 4.5. táblázat mutatja a teljes tétellistát átlagértékek szerint rendezve.

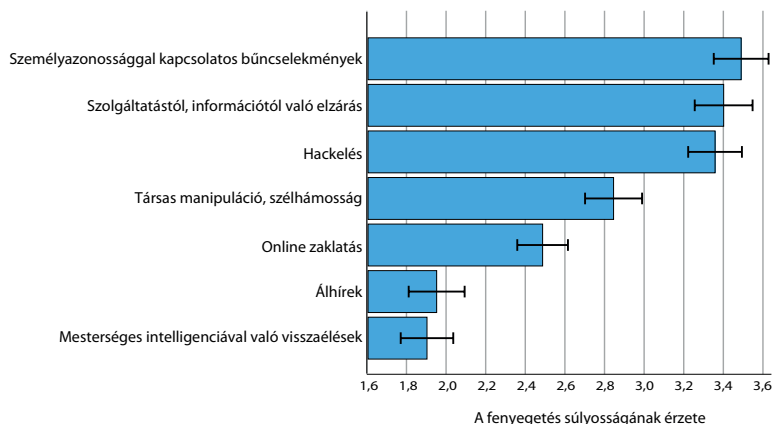
*4.5. táblázat: Az online visszaélések súlyosságának
megítélése – a fenyegetés súlyosságának érzetét mérő tételek
leíró statisztikái az átlag szerint csökkenő rendben*

	Átlag	Szórás
Valaki megszerzi a banki adataimat, és a nevemben vásárol, átutal, hitelkártyát vált ki vagy kölcsönt vesz fel.	3,93	1,190
Egy zsarolóvírus titkosítja és hozzáférhetetlenné teszi a számítógépemen tárolt adatokat, és az elkövető váltságdíjat követel a titkosítás feloldásáért.	3,55	1,198
A számítógépemre vagy okostelefonomra olyan kémprogram kerül, amely a tudtom nélkül felvételeket készít rólam a saját eszközőm kamerájával.	3,53	1,226
A számítógépemre vagy okostelefonomra olyan kémprogram kerül, amely megfigyeli és illetéktelen személynek közvetíti a képernyőn látható információkat vagy a billentyűleütéseket.	3,47	1,151
Nyilvános helyen valaki – személyesen vagy kamera használatával – elvesz egy jelszót, amelyet begépelek, így hozzáfér bizonyos adataimhoz.	3,36	1,168
A számítógépemre vagy okostelefonomra kémprogram kerül, amely segítségével illetéktelen személyek hozzáférhetnek az eszközön tárolt adatokhoz.	3,36	1,183
Valaki – a rólam online elérhető információk alapján – egyszerű próbálkozással, találgatással kitalálja az egyik jelszavam, így hozzáfér bizonyos adataimhoz.	3,26	1,158
Olyan hackertámadás ér egy általam is igénybe vett szolgáltatást (például bank, tanulmányi rendszer, ügyfélkapu), amely számomra is bosszúságot vagy kárt okoz.	3,24	1,194
A számítógépemre vagy okostelefonomra vírus vagy más rosszindulatú szoftver kerül.	3,14	1,176
Egy haragosom nyilvánosan közzéteszi a személyes adataimat (lakcímem, telefonszámom, e-mail-címem stb.), hogy így álljon rajtam bosszút.	3,04	1,149
Egy ismeretlen személy érzékeny információkat kezd gyűjteni rólam, követi az online tevékenységemet, és mindezt a tudomásomra is hozza, hogy ezáltal félelmet keltsen bennem, és ne érezzem magam biztonságban.	2,94	1,242
Valaki azzal próbál zsarolni, hogy intim felvételekkel rendelkezik rólam, és ha nem teljesítem a feltételeit, ezeket eljuttatja rokonaimnak, barátaimnak és kollégáimnak.	2,90	1,376

	Átlag	Szórás
Egy csaló megtéveszt azzal, hogy valamilyen szolgáltatónak (banknak, közüzemnek, futárszolgálatnak) adja ki magát, és ráveszt, hogy megadja magamról valamilyen személyes adatot.	2,86	1,177
Egy csaló megtéveszt azzal, hogy személyes ismerősömnek vagy ismerősöm barátjának adja ki magát, és így vesz rá valamire.	2,82	1,151
Tanúja leszek annak, hogy egy online üzenetben valaki egy olyan társadalmi csoport ellen buzdít erőszakra, amelynek én is tagja vagyok (például származásom, szexuális orientáció, vallásom vagy politikai szimpátiám alapján).	2,45	1,170
Egy ismeretlen személy rágalmozó üzeneteket küld rólam az ismerőseimnek.	2,40	1,143
Valakitől ismétlődően bántó, sértő tartalmú üzeneteket kapok, amelyek a küllememet, a tulajdonságaimat vagy a képességeimet/hozzáértésemet becsmérlik.	2,16	1,133
Egy online közösségben, amelynek tagja vagyok, valaki provokatív megjegyzéseivel szándékosan feszültséget gerjeszt, és engem is támadni kezd (trollkodás).	2,04	1,052
Egy mesterséges intelligencia által generált, személyesen nekem címzett szöveges üzenetről elhiszem, hogy azt ember (például ügyfélszolgálati munkatárs) fogalmazta.	2,00	1,081
Egy hivatalos hírmédia által közzétett hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.	1,99	1,125
Egy közösségi médiában terjedő, bizonytalan eredetű hír tartalmát elhiszem, de később kiderül, hogy az álhír volt.	1,91	1,068
Egy mesterséges intelligencia által generált fényképről elhiszem, hogy az valódi.	1,80	1,016

Forrás: saját szerkesztés

Jól látható, hogy bár az álhírekkel és az MI-tartalmakkal való sikeres megtévesztést valószínű eseménynek tartják, valójában ezek bekövetkeztétől tartanak az adatközlők a legkevésbé. Az életüket legsúlyosabban érintő problémák közé a személyes adataikat és vagyonukat közvetlenül veszélyeztető visszaéléstípusokat sorolták. Ezt támasztja alá az összesített mutatók rangsorolása is, amely szerint például a hackertámadásokat súlyos fenyegetésnek érzik, azzal együtt, hogy – mint fent láttuk – igen valószínűnek is tartják, hogy a jövőben ilyen támadások áldozatává válhatnak (4.7. ábra).



4.7. ábra: A fenyegetés súlyosságának érzete visszaéléstípusok szerint összesítve (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

Válaszhatékonyság

Szintén érdekes tendenciák figyelhetők meg, ha összevetjük, mely védekezési módok hatékonyságában hisznek az adatközlők a leginkább, illetve legkevésbé. Az átlag szerint csökkenő sorba rendezett itemlistát a 4.6. táblázat tartalmazza.

4.6. táblázat: A különféle védekezési módok hatékonyságának megítélése – a válaszhatékonyaságot mérő tételek leíró statisztikái az átlag szerint csökkenő rendben

	Átlag	Szórás
Ne osszunk meg magunkról nyilvánosan érzékeny, személyes információkat!	4,20	0,971
A jelszavainkat rendszeresen változtassuk meg!	3,99	0,967
Minden platformon más, egyedi jelszót használjunk!	3,95	1,004
Banki vagy más közületi alkalmazásból/honlapról az ügyintézés végén azonnal jelentkezzünk ki!	3,94	,986
Senkinek ne engedjük, hogy olyan fénykép- vagy videófelvételeket készítsen rólunk, amelyekkel később bárki más visszaélhet!	3,94	1,103
Ne használjunk olyan jelszót, amely családtagjaink/házikedvencünk nevét, monogramját, születési évszámot, hobbinkkal kapcsolatos szót vagy más könnyen kitalálható információt tartalmaz!	3,92	1,029
Videóhívás közben ne tegyünk olyat, amiről ránk nézve kínos vagy kétértelmű képernyőkép készíthető!	3,86	1,078
Online vásárlásra tartunk fenn külön számlát és bankkártyát, amelyen egyébként nem tartunk jelentős összeget!	3,84	1,104
Privát közösségimédia-profilát használunk, amelyen kizárólag a beleegyezésünkkel válhat valaki a követőnkké, ismerősünkké!	3,80	1,082
Ne látogassunk olyan weboldalakot, amelyek ingyen kínálnak egyébként fizetős szolgáltatást!	3,80	1,100
Személyes adatainkat (e-mail-cím, telefonszám) ne tegyük közzé nyilvános felületen!	3,78	1,141
Ne bízzuk magunkat az operációs rendszer beépített vírusvédelmére: vegyünk igénybe professzionális vírusvédelmet adó szolgáltatást!	3,68	1,061
Ha nyilvános helyen kell egy fontos jelszót begépelnünk, győződjünk meg róla, hogy nincs senki a közelünkben, és takarjuk el a képernyőt/billentyűzetet, amennyire csak lehetséges!	3,67	1,067
Ha egy hivatkozás intézményi weboldalra irányít, mindig ellenőrizzük az oldal URL-címét!	3,67	1,054
Rendszeresen frissítsuk a számítógépünk/okostelefonunk operációs rendszerét és a használt alkalmazásokat!	3,67	1,168
Takarjuk le a számítógépünk kameráját, ha az nincs használatban!	3,66	1,192
Ha valaki sértő vagy gyűlölködő tartalmú üzeneteket küld, szakítsuk meg vele az online kapcsolatot (töröljük, tiltsuk le)!	3,64	1,222

	Átlag	Szórás
Ha régi vagy távoli ismerős veszi fel velünk a kapcsolatot, igyekezzünk ellenőrizni, hogy a személy valóban az-e, akinek mondja magát!	3,62	1,056
Az online olvasott hírek tartalmát mindig ellenőrizzük más forrásokban is!	3,62	1,141
A valószínűtlen helyzeteket / eseményeket ábrázoló fényképeken ellenőrizzük az apró részleteket, amelyek esetleg arról árulkodnak, hogy a felvétel nem valódi!	3,60	1,092
Intézménytől érkező üzenetek esetében mindig ellenőrizzük a feladót!	3,56	1,104
Ne töltsünk le jogvédett szoftvereket fájlmegosztó rendszerekről!	3,54	1,105
Ne használjunk olyan szolgáltatást, amely a jelszavaink elmentésével megkönnyíti a különféle fiókokba való belépést!	3,46	1,163
Ne töltsünk le egyéb jogvédett tartalmakat (például filmek, könyvek) fájlmegosztó rendszerekről!	3,33	1,143
Online közösségekben, közösségi médiában ne fogalmazzunk meg olyan álláspontot, amely miatt gyűlölködés célpontjává válhatunk!	3,25	1,351
Rendszeresen töröljük a böngészőnk által eltárolt sütitket!	3,23	1,165
A fontos, határidős online ügyleteket (például banki átutalás, adóbevallás) ne halogassuk az utolsó pillanatra!	3,16	1,349
A nem garantáltan megbízható weboldalak meglátogatásakor használjuk a böngészőnket inkognitó módban!	3,10	1,295
A számítógépünkön ne rejtjük el az állományok kiterjesztését (például „pdf”, „docx”, „docm”, „exe”)!	3,05	1,235

Forrás: saját szerkesztés

A hiedelmeket összesítve tükröző értékek tehát azt mutatják, hogy a kitöltők leginkább jelszavak gondos kezelését és a tartalommegosztással kapcsolatos tudatosságot tekintik hatékony védekezési eszközöknek, a legkisebb jelentőséget pedig olyan intelmeknek tulajdonítanak, amelyek szerint biztonsági kockázatot jelenthet és kerülendő a nem hivatalos forrásból származó állományok letöltése, a bizonytalan hátterű weboldalakról a számítógépünkre kerülő süтик, az online ügyintézés közbeni kapkodás, vagy az extrém véleménykifejezés.

Önhatékonyság

Végül érdekes mintázatok figyelhetők meg a kitöltők saját képességeire vonatkozó kérdések esetében is. Átlag szerint csökkenő rendben az önhatékonyságra vonatkozó tételek rangsorát a 4.7. táblázat mutatja.

4.7. táblázat: A saját képességekre vonatkozó hiedelmek – az önhatékonyságot mérő tételek leíró statisztikái az átlag szerint csökkenő rendben (minden állítás a „Képes vagyok...” mondatkezdetet egészíti ki)

	Átlag	Szórás
...figyelni arra, hogy ne osszak meg magamról online érzékeny, személyes információkat.	4,32	0,851
...arra, hogy a fontos felhasználói vagy ügyfélprofiljaimat biztonságos jelszóval védjem meg.	4,20	0,917
...vigyázni arra, hogy nyilvános helyen senki ne lássa az általam begépett jelszót/PIN-kódot.	4,16	0,939
...a biztonságom érdekében átgondoltan testre szabni a közösségimédia-profilom beállításait.	4,14	0,918
...garantálni az online vásárlásaim és banki adataim biztonságát.	4,12	0,914
...a jelszavaimat biztonságosan kezelni.	4,08	0,916
...figyelni arra, hogy kizárólag olyan üzeneteket írjak online, amelyeket a későbbiekben is vállalok, és nem érintene károsan, ha nyilvánosságra kerülnének.	4,04	0,983
...ellenőrizni, hogy egy intézmény valódi weboldalát látogattam meg, és nem annak hamisított másolatát.	3,97	0,949
...biztonságosan használni a böngészőprogramot.	3,95	0,874
...megkülönböztetni a biztonságos weboldalakat a gyanús és potenciálisan veszélyes honlapoktól.	3,95	0,871
...ellenőrizni, hogy egy üzenet hiteles forrásból származik-e.	3,92	0,880
...arra, hogy megkülönböztessem az álhíreket a valódiaktól.	3,87	0,887
...figyelni arra, hogy semmilyen körülmények között ne készülhessen rólam olyan fénykép- vagy videófelvétel, amelyről nem szeretném, ha nyilvánosságra kerülne.	3,80	1,038

	Átlag	Szórás
...arra, hogy odafigyeljek az általam használt operációs rendszerek és egyéb szoftverek rendszeres frissítésére.	3,79	1,084
...arra, hogy ne halogassam a határidős, online eszközöket igénylő feladatokat az utolsó pillanatra.	3,73	1,085
...vírusirtó és egyéb biztonsági szoftverek telepítésére, használatára és rendszeres frissítésére.	3,67	1,117
...a számítógépen megkülönböztetni a futtatható állományokat az adatokat tartalmazó állományoktól.	3,60	1,110
...arra, hogy megkülönböztessenem a mesterséges intelligencia által generált szövegeket az ember által fogalmazott szövegektől.	3,55	0,974
...minden platformon eltérő jelszót használni.	3,55	1,167
...arra, hogy a számomra szükséges jogvédett tartalmakat (például könyvek, filmek) mindig jogszerű forrásból szerezzem be, és ellenálljak az ingyenes, de nem jogszerű források kísértésének.	3,52	1,207
...arra, hogy megkülönböztessenem a mesterséges intelligencia által generált fényképeket a valódi fényképfelvételektől.	3,50	1,004

Forrás: saját szerkesztés

Látható például, hogy míg az adatközlők képesnek érzik magukat arra, hogy biztonságos jelszóval védjék meg elektronikus fiókjaikat, úgy érzik, ahhoz nem rendelkeznek elegendő önfegyelmekkel, hogy az egyes platformokon különböző jelszót használjanak. További érdekes eredmény, hogy az ingyenesen elérhető, nem jogtisztta forrásból származó tartalmak kísértésének való ellenállás különösen nagy kihívást jelent annak ellenére, hogy az ilyen tartalmak rendszeres letöltése nyilvánvaló kiberbiztonsági kockázattal jár.

A védelemmotiváció összetevői közötti korrelációk

A kutatósmódszertani szak- és tankönyvirodalomban jól ismert intelem, hogy a korrelációs kapcsolatok nem jelentenek bizonyítékot oksági összefüggésekre. Ez természetesen helyes és fontos kitétel, hiszen két változó között akár erős korrelációt is eredményezhet egy harmadik változó – egy

ügynevezett közös ok – hatása is. Mindezzel együtt azonban azt is fontos látnunk, hogy amennyiben változók egy csoportján belül meghatározott rendben érvényesülnek oksági hatások, akkor ennek az összefüggésrendszernek gyakran felismerhető lenyomata lesz a változók közötti kapcsolatok erősségét és irányát páronként jellemző korrelációs mátrixban. Tehát bár a korrelációs kapcsolatok szigorú értelemben sosem bizonyítanak oksági összefüggéseket, ha már rendelkezünk egy bizonyos elképzeléssel egy mögöttes oksági modell viszonyrendszeréről, ellenőrizhetjük azt, hogy a korrelációs mátrix az elképzelésünkkel összhangban van-e vagy annak ellentmondó mintázatot mutat. Ez a fajta összevetés elvégezhető olyan modern módszerekkel, mint a strukturáltegyenlet-modellezés, vagy hagyományosabb megközelítésekkel, mint az egyszerű többváltozós regressziókon alapuló útelemzés. (A területet és az egyes megközelítések előnyeit és hátrányait részletesen áttekinti Keith).⁵⁰

Kutatásunk leíró-feltáró jellegéből következően nem rendelkezünk előzetes modellel vagy hipotézisekkel a kiberbiztonsági védelemmotiváció komponensei közötti oksági összefüggésekről. Mindazonáltal érdemes megjegyezni, hogy a négy összetevő közötti kapcsolatokat leíró korrelációs mátrix (4.8. táblázat) olyan tiszta mintázatot mutat, amely jól értelmezhető egy logikus és egyben egyszerű oksági lánc megnyilvánulásaként.

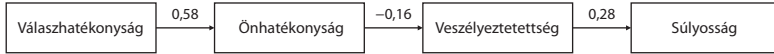
4.8. táblázat: A védelemmotiváció négy dimenziója közötti korrelációk

	1.	2.	3.
1. Személyes veszélyeztetettség			
2. A fenyegetés súlyossága	0,282**		
3. Válaszhatékonyosság	0,003	0,054	
4. Önhatékonyosság	-0,159*	-0,008	0,583**

Megjegyzés: * $p < 0,05$; ** $p < 0,01$

Forrás: saját szerkesztés

⁵⁰ KEITH 2019.



4.8. ábra: A védelemmotiváció oksági modellje (a feltüntetett számérték minden esetben a szignifikáns korreláció Pearson-féle együtthatója)

Forrás: saját szerkesztés

Egy oksági láncban természetes, hogy két szomszédos változó között – vagyis mikor az oksági hatás közvetlen módon érvényesül – erősebb korrelációt találunk, mint két olyan változó között, amelyek az oksági hatások sorában távolabb helyezkednek el egymástól, hiszen ez utóbbi esetben a hatások közvetettek, és azok érvényesülését egyéb, nem vizsgált hatások és a mérési hibából eredő zaj részben elfedheti.

Esetünkben megfigyelhető, a négy tényező sorba rendezhető oly módon, hogy kizárólag a láncolat szomszédos tagjai korrelálnak egymással szignifikánsan, az ennél távolabb álló tagok között az összefüggés már nem szignifikáns. A korrelációs mintázat tehát összhangban áll a 4.8. ábrán látható kauzális modell struktúrájával.

A fenti oksági sor nemcsak összhangban van az általunk feltárt korrelációs struktúrával, hanem ezzel együtt egy igen logikus és könnyen értelmezhető narratívának is megfelel. Első lépésben az egyén meggyőződik arról, hogy igenis lehetséges az online világban ránk leselkedő fenyegetések ellen hatékonyan védekezni (válaszhatékonyság). Ennek következményeképpen – miután megismerkedett a hatékony védekezés módozataival – meggyőződik arról, hogy a védekezési stratégiák nem igényelnek különösebb szakmai képzettséget, megfelelő odafigyeléssel és önfegyelemmel saját maga is képes azokat végrehajtani (önhatékonyság). Ha pedig az egyén már hisz abban, hogy ő maga képes hatékonyan védekezni, kevésbé fogja valószínűnek tartani, hogy online visszaélés áldozatául essen (negatív korreláció a veszélyeztetettséggel). Végül az, aki kevésbé érzi magát veszélyeztetve, kevésbé is érzi súlyos problémának a kiberbiztonsági fenyegetéseket.

Még egyszer fontos hangsúlyozni, hogy az adatsorunk nem jelent közvetlen bizonyítékot a fenti oksági láncra: ezen a ponton pusztán annyi állapítható meg, hogy a feltárt korrelációs kapcsolatok mintázata összhangban áll egy ilyen felépítésű oksági modellel.

ÖSSZEGZÉS

A PMT alapelve, hogy a védelemmotivációt a fenyegetés észlelése (fenyegetettség, súlyosság) és a megküzdési értékelés (válaszhatékonyság, önhatékonyság) határozza meg. Az elmélet szerint a magas fenyegetettség és az eredményes védekezési lehetőségek észlelése növeli az egyének motivációját a preventív viselkedésre.

A védelemmotiváció elméletére épülő kutatás célja az volt, hogy feltárja a digitális fenyegetések észlelésének, valamint a preventív magatartások alkalmazásának pszichológiai összefüggéseit. Az elmélet négy kulcsfontosságú tényezőt különít el: a fenyegetettség érzékelését, a fenyegetések súlyosságának megítélését, a védekezési stratégiák hatékonyságába vetett hitet (válaszhatékonyság), valamint az egyének saját védekezési képességeikbe vetett bizalmát (önhatékonyság). A kutatás során kérdőíves adatgyűjtést alkalmaztunk, amelyben a résztvevők különféle digitális fenyegetésekkel kapcsolatos érzékelésükről, azok megítélt súlyosságáról és a védekezési módokról alkotott véleményükről számoltak be. A kérdőív tételei kiberbűnözési kategóriák szerint strukturálták a vizsgált fenyegetéseket, beleértve a hackertámadásokat, az online zaklatást, az adathalászatot, az álhíreket és a mesterséges intelligenciával történő visszaéléseket. A kutatás kiemelt figyelmet fordított a fenyegetések valószínűségének és súlyosságának mérésére, amely Likert-skálás értékelésekkel történt.

Az eredmények számos fontos mintázatra világítottak rá. A fenyegetettség érzékelése alapján a résztvevők leginkább a mesterséges intelligenciával előállított tartalmak, például hamisított szövegek és képek által okozott megtévesztések valószínűségét ítélték magasnak. Emellett jelentős fenyegetést láttak a hackertámadásokban, például rosszindulatú szoftverek vagy kémprogramok telepítésében, illetve szolgáltatások elleni támadásokban. Az álhírek is előkelő helyet foglaltak el a megítélt fenyegetések listáján, különösen azok, amelyek hivatalosnak tűnő forrásokból származnak. Ezzel szemben a résztvevők kevésbé tartották valószínűnek, hogy szélhámosok vagy online csalók sikeresen megtévesztik őket személyes adatok megszerzésére vagy cselekvésre való ráhatás céljából. A megítélt fenyegetettség tehát nagyban függött a fenyegetés típusától és a résztvevők

egyéni tapasztalatától, valamint attól, hogy mennyire tartották valószínűnek az esemény bekövetkezését.

A fenyegetés súlyosságát illetően a pénzügyi és adatbiztonsági visszaélések emelkedtek ki a legkritikusabb problémaként. A résztvevők szerint a banki adatokkal való visszaélés, például hitelkártyacsалások vagy a zsarolóvírusokkal való támadások, rendkívül súlyos következményekkel járhatnak. Ezeket követték azok az esetek, amelyekben érzékeny adatok kerülnek illetéktelen kezekbe, például kémprogramok vagy személyes információk online nyilvánosságra hozatala révén. Érdekes módon, bár a mesterséges intelligenciával előállított tartalmakkal kapcsolatos fenyegetések valószínűségét magasnak ítélték, ezek súlyosságát alacsonynak értékelték. Ez arra utalhat, hogy az ilyen típusú megtévesztések kevésbé érintik közvetlenül a résztvevők anyagi vagy pszichológiai biztonságát, mint például a hackertámadások vagy a pénzügyi adatokat érintő visszaélések.

A válaszhatékony-ság dimenziójában a résztvevők leginkább a jelszavak gondos kezelését, valamint az érzékeny információk nyilvános megosztásának kerülését tartották hatékony preventív stratégiának. Kiemelték továbbá, hogy a privát közösségimédia-profil használata, a számítógépes rendszerek rendszeres frissítése és a biztonsági szoftverek alkalmazása is hozzájárulhat a fenyegetések megelőzéséhez. Az önhatékony-ság mérésében az egyének saját képességeiket illetően általánosan magas értékeket mutattak, különösen olyan stratégiák esetében, mint a hiteles források ellenőrzése, a jelszavak biztonságos kezelése vagy a biztonságos online vásárlási gyakorlatok alkalmazása. Ugyanakkor a résztvevők kisebb mértékű önhatékony-ságot jeleztek olyan helyzetekben, amelyek fokozott önfegyelmet igényeltek, például amikor minden platformon különböző jelszavak használatára vagy a nem jogtisztá tartalmak elkerülésére volt szükség.

Az összefüggések elemzése megmutatta, hogy a védelemmotiváció egyes összetevői szignifikánsan korrelálnak egymással. A válaszhatékony-ság növelése közvetlenül hozzájárult az önhatékony-ság erősödéséhez, ami viszont csökkentette a fenyegetettség érzését. Ez az oksági lánc arra utal, hogy ha az egyének hatékony-nak tartják a rendelkezésre álló védekezési stratégiákat, nagyobb valószínűséggel érzik úgy, hogy képesek is alkalmazni azokat, ami csökkenti a fenyegetések iránti érzékenysé-güket. Az eredmények

szerint a fenyegetettségérzet csökkenésével párhuzamosan a fenyegetések súlyosságának érzékelése is mérséklődött, ami tovább erősítette a preventív magatartás elfogadását.

Összességében a kutatás eredményei alátámasztják, hogy a PMT hatékony keretet nyújt az egyének biztonságtudatos viselkedésének elemzéséhez. Az eredmények gyakorlati jelentőségűek a biztonságtudatosságot növelő kampányok tervezésében, különösen azokban, amelyek a válaszhatékony-ságot és az önhatékony-ságot célozzák. A kutatás eredményei szerint a digitális fenyegetésekkel szembeni védekezésben a személyes tapasztalatok, a fenyegetések típusa és a védekezési lehetőségek hatékonyságába vetett bizalom meghatározó szerepet játszanak. Ezek az eredmények fontos alapot szolgáltatnak a további kutatásokhoz és a kiberbiztonsági stratégiák fejlesztéséhez.

A kutatás eredményei rámutattak, hogy a PMT keretei között végzett vizsgálatok hatékonyan feltárják az egyének biztonságtudatos viselkedésének indikátorait. Az eredmények gyakorlati jelentőségűek, különösen a biztonságtudatosságot növelő kampányok tervezése és végrehajtása szempontjából, amelyek a válaszhatékony-ságot és az önhatékony-ságot célozzák meg.

5. A DEMOGRÁFIAI CSOPORTOK KÖZÖTTI KÜLÖNBSÉGEK A VÉDELEMMOTIVÁCIÓBAN

A kiberbiztonság területén az emberi tényezők figyelembevétele nélkülözhetetlen, különösen a demográfiai különbségek szempontjából, amelyek jelentős hatással lehetnek a védelemmotivációra. Az emberi viselkedés és a technológiai megoldások közötti interakciók alapos megértése elősegíti az egyéni és szervezeti szintű biztonság fokozását. A kutatások kiemelik, hogy az online fenyegetésekre adott válaszokat olyan tényezők befolyásolják, mint az életkor, a nem, az iskolai végzettség, valamint az általános technológiai tapasztalatok és attitűdök.

A demográfiai különbségek feltérképezése tehát kulcsfontosságú az online fenyegetésekkel szembeni viselkedés megértéséhez, mivel az egyes csoportok eltérően érzékelhetik az online visszaélések jelentette fenyegetést, és másként értékelhetik azok potenciális hatásait. Az is különböző mértékben jelenhet meg, hogy hisznek-e az eredményes védekezés lehetőségében, illetve hogy mennyire látják magukat képesnek a szükséges intézkedések megtételére. Ezek az attitűdök alapvetően befolyásolják a kiberbiztonsági stratégiák és képzési programok hatékonyságát. Az életkor vonatkozásában a nemzetközi szakirodalmi álláspont az, hogy a fiatalabb generációk technológiai önhatékonysága magasabb, ami gyakran alacsonyabb fenyegetettségérzethez vezet, így kevésbé hajlandók védelmi intézkedéseket elfogadni. Ezzel szemben az idősebb korosztály gyakran érzékenyebb a kiberfenyegetésekre, de alacsonyabb technológiai önbizalommal rendelkezik, ezért hatékony oktatás és egyszerű megoldások szükségesek számukra. A kutatások szerint az életkor tehát közvetlen hatással van a fenyegetés észlelésére és a védelmi viselkedésre. A nem tekintetében a közöttük fennálló különbségek szintén jelentősek: a férfiak gyakrabban rendelkeznek magasabb technológiai önhatékonysággal, de hajlamosak alábecsülni a fenyegetéseket. A nők ezzel szemben nagyobb sebezhetőséget érzlelnek, ami növeli védelmi viselkedésük valószínűségét, ha hatékony és praktikus megoldásokat kínálnak nekik.

Az oktatási programok nemi érzékenysége pedig képes ezáltal növelni az információbiztonsági tudatosságot. Az iskolai végzettség tekintetében a magasabb iskolai végzettséggel rendelkezők általában jobban megértik a fenyegetések súlyosságát és a védelmi intézkedések hatékonyságát, ezért hajlamosabbak a tudatosabb viselkedésre. Az alacsonyabb iskolai végzettségűek számára egyszerű, vizuális és könnyen hozzáférhető oktatási anyagok szükségesek, amelyek növelik az önhatékonyt.⁵¹

Az Egyesült Államokban működő Capitol Technology University és az Egyesült Királyságban működő Edge Hill University kutatóinak álláspontja a fent megfogalmazottakkal korrelál: az online fenyegetésekkel szembeni motivációban megfigyelhető demográfiai különbségek közül kiemelkedőek a nemi alapú eltérések. Ezen kutatás szerint a nők hajlamosabbak a fenyegetettség súlyosságát erőteljesebben érzékelni, míg a férfiak gyakran alábecsülik a kockázatokat, ami eltérő védelemi stratégiákhoz vezethet. Az iskolai végzettség is fontos tényező: a magasabb végzettségűek általában nagyobb önhatékonyt mutatnak, és hajlamosabbak a tudatos védekezési technikák alkalmazására. Egy másik kulcsfontosságú megállapítás az, hogy az emberközpontú megközelítések – mint például a viselkedési profilalkotás – hatékony eszközök lehetnek a kiberbiztonsági kihívások kezelésében. Ezek a módszerek az egyének viselkedési mintáinak elemzésére épülnek, és segítenek azonosítani a fenyegetések forrásait, valamint a hatékonyabb megelőzési stratégiák kialakítását.⁵²

A nemzetközi szakirodalom alapján a kiberbiztonsági kihívásokra adott megoldások nem merülhetnek ki csupán technológiai eszközök alkalmazásában, mivel a felhasználók viselkedése, döntései és attitűdjei alapvetően befolyásolják ezek hatékonyságát. Az emberi hibák – akár figyelmetlenség, akár tapasztalatlanság révén – a biztonsági rendszerek leggyakoribb gyenge pontjai, amelyeket célzott képzésekkel és tudatosítással lehet csökkenteni, ezen kampányok sikerességének pedig elengedhetetlen elemei a demográfiai tényezők.⁵³ Az emberi tényezőkből adódó változók,

⁵¹ KHADER–KARAM–FARES 2021; SNIDER et al. 2021; ALMANSRI – AL-EMRAN – SHAALAN 2023.

⁵² MARTINEAU–SPIRIDON–AIKEN 2023.

⁵³ ROHAN et al. 2021.

például a *multitasking* vagy az információs túlterheltség és a nem megfelelő képzés szintén megkönnyítik a humán alapú sérülékenységek kihasználását, így az oktatás és a képzés központi szerepet játszanak mindezek enyhítésében.⁵⁴

Összességében megállapítható, hogy az egyének tudatossága és attitűdje a kiberfenyegetések súlyosságáról, valamint védekezési képességeikről jelentősen eltér demográfiai csoportonként. A biztonságtudatos magatartás fejlesztése célzott oktatási programok révén érhető el, amelyek figyelembe veszik ezeket az eltéréseket.⁵⁵

A jövőbeni kutatások célja, hogy még mélyebben megértsék a demográfiai tényezők és az online viselkedés közötti kapcsolatot, és olyan oktatási programokat fejlesszenek ki, amelyek a különböző társadalmi csoportok specifikus igényeit célozzák meg.⁵⁶

Mindezekből adódóan a következő fejezetben részletesen bemutatjuk a különböző demográfiai tényezők és a védelemmotiváció közötti összefüggéseket az alábbi témakörökben:

- Nemek közötti különbségek: A nők és férfiak eltérően érzékelik az online fenyegetések súlyosságát, és másképp reagálnak a védekezés szükségességére. A kutatás részletezi a nemekre vonatkozó szignifikáns eredményeket, például a fenyegetettség érzékelt súlyosságában és a védelemhez való hozzáállásban megfigyelhető különbségeket.
- Iskolai végzettség hatása: Az oktatás szintje és a technológiai önhatékonyság közötti összefüggést mutatjuk be, különös figyelemmel arra, hogyan értékelik saját védekezési képességeiket a magasabb iskolai végzettségű csoportok. A fejezet részletezi az önhatékonyság mérésére használt elemzési módszereket és az ezzel kapcsolatos megállapításokat.
- Aktív kereső státusz: A fejezet vizsgálja az aktív keresők és a nem keresők közötti különbségeket az önhatékonyság és a védelemmotiváció többi komponense kapcsán.

⁵⁴ MARTINEAU–SPIRIDON–AIKEN 2023.

⁵⁵ JEONG et al. 2019.

⁵⁶ ROHAN et al. 2021.

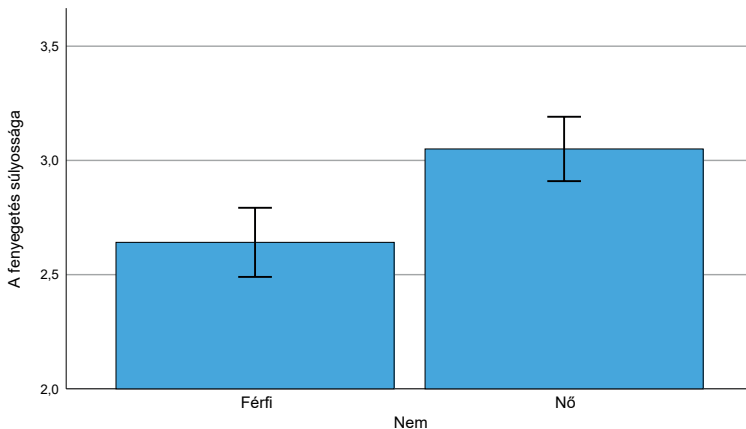
- ♦ Kiberbiztonsági képzések hatása: Részletezzük a kiberbiztonsági képzések hatásait, különösen az önhatékonyság növelésére és a védekezési stratégiák javítására vonatkozóan. A kutatás bemutatja, hogy a képzések miként befolyásolják a fenyegetettség érzékelését és a motivációt.

EREDMÉNYEK

Az 1. és 4. fejezetekben a demográfiai változók és a védelemmotiváció tényezőinek módszertani kérdéseit részletesen tárgyaltuk, így ebben a fejezetben az eredmények közlésére szorítkozunk. A 3. fejezet felépítéséhez hasonlóan előbb a szignifikáns eredményeket adó próbák eredményeit közöljük, és a későbbiekben összesítve említjük meg mindazokat az elemzéseket, amelyek nem mutattak ki összefüggést a demográfiai változók és a védelemmotiváció elemei között.

Nemek közötti különbségek a fenyegetés észlelt súlyosságában

A nemek és a védelemmotiváció négy összetevője közötti kapcsolatát négy független mintás t -próbával ellenőriztük, és ezek közül egy esetben jelzett a próba szignifikáns eltérést a két nem között: a nők súlyosabbnak látják az online visszaélések következményeit, mint a férfiak, azaz hajlamosabbak úgy megítélni, hogy ha ők ilyen típusú visszaélés áldozatává válnának, az rendkívül súlyos problémát jelentene számukra ($t_{[213]} = -3,847; p < 0,001$). E közepes mértékű hatást (Cohen-féle $d = -0,528$) az 5.1. ábra szemlélteti, jól látható, hogy a konfidenciaintervallumok nem fednek át egymással.

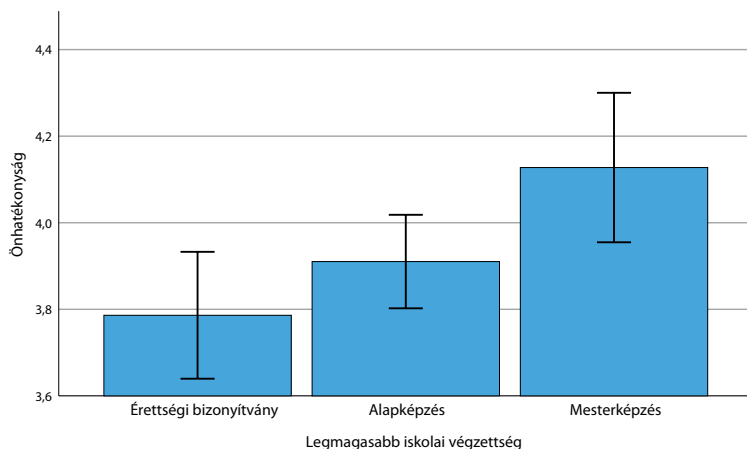


5.1. ábra: A nők szignifikánsan súlyosabbnak ítélik meg az online visszaélések következményeit, mint a férfiak (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

Iskolai végzettség és önhatékonyság

A születési nemhez hasonlóan az iskolai végzettség is csak egy esetben mutatott szignifikáns összefüggést a védelemmotivációval, ez esetben azonban az önhatékonyság bizonyult a demográfiai változóval szignifikánsan összefüggő komponensnek. Az egyszempontos varianciaanalízis ($F[2, 208] = 3,832$; $p = 0,023$) és a Bonferroni-korrekcióval végzett *post hoc* elemzés eredményei szerint a mesterképzésen szerzett oklevéllel rendelkezők sokkal inkább gondolják úgy, hogy képesek hatékonyan védekezni az online fenyegetések ellen, mint azok, akik (még) nem rendelkeznek diplomával. Ahogy az 5.2. ábrán is látható, míg az alapképzésen szerzett diplomával rendelkezők átlagos önhatékonysága a mintánkban e két csoport közötti értéket vett fel, a konfidenciaintervallumok közötti átfedés a másik két csoporttal esetükben túl nagymértékű ahhoz, hogy bármely csoporttól való eltérésüket bizonyítottan tekintsük.



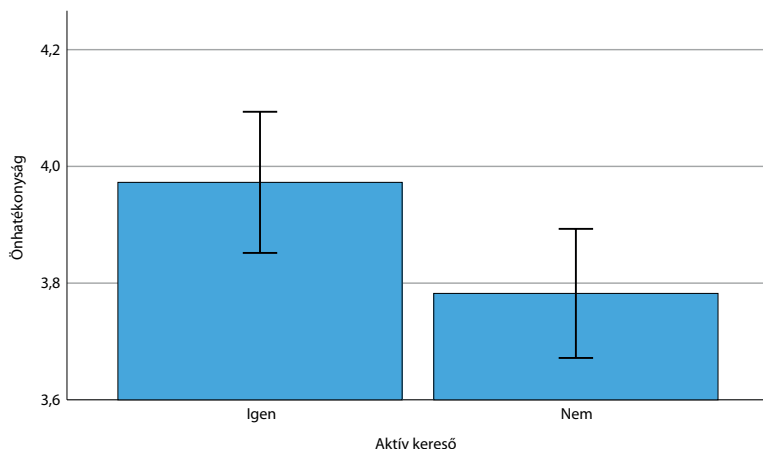
5.2. ábra: Az iskolai végzettség és az önhatékony-ság összefüggése
(a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

Az ábrán látható mintázat igen hasonló ahhoz, mint amelyet a végzettség és az aktív védekezés A2 (közzétett tartalmak körének és terjedésének kontrollja) dimenziójának esetében láttunk: bár a lépcsőzetes mintázat a mintaátlagokon megfigyelhető, a jelen mintanagyság és varianciák mellett a *post hoc* elemzésnek nincsen olyan statisztikai ereje, hogy ilyen nagyságrendű hatás (ha létezik is) egyértelműen igazolható legyen. Tehát bár az ANOVA alapján kijelenthetjük, hogy a végzettség és az önhatékony-ság közötti kapcsolatra az adatsorunk bizonyítékot jelent, az nem látjuk biztosan, hogy a magasabban képzettek előnye már az alapszintű oklevél megszerzésével megjelenik-e, vagy a mesterszintű diploma megszerzésével áll kapcsolatban (vagy két lépésben, lépcsőzetesen érvényesül).

Az aktív kereső státusz és az önhatékonyság

A végzettséghez hasonlóan az aktív kereső státusz is kizárólag az önhatékonysággal áll szignifikáns összefüggésben, a védelemmotiváció többi komponense tekintetében nem detektálható eltérés a keresők és nem keresők között. Az eredmények szerint az aktív keresők jobban bíznak abban, hogy képesek magukat megvédeni az online fenyegetésektől, mint aki egyelőre nem rendelkezik rendszeres jövedelemmel ($t[213] = 2,301, p = 0,022$). Ahogy az 5.3. ábrán látható, az eltérés kismértékű (Cohen-féle $d = 0,314$), de már szignifikáns:



5.3. ábra: Az aktív keresők hajlamosabbak úgy gondolni, hogy képesek hatékonyan védekezni az online visszaélések ellen, mint a rendszeres jövedelemmel nem rendelkezők (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

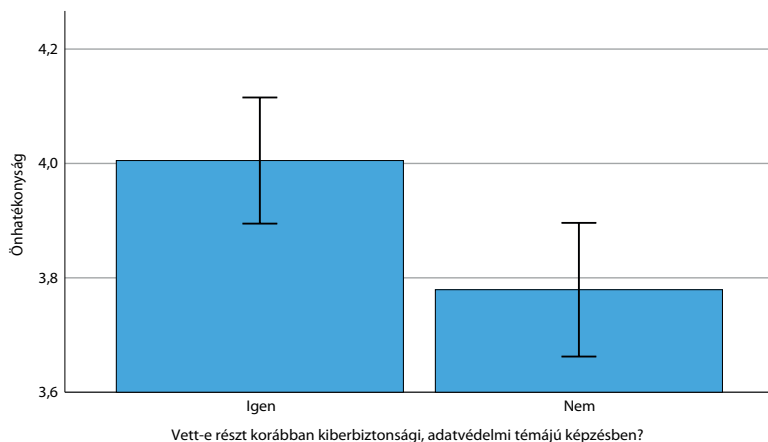
Forrás: saját szerkesztés

A kiberbiztonsági képzések hatása a védelemmotivációra

Kutatásunk szempontjából fontos kérdés volt, hogy feltárjuk, milyen kimutatható hatásai vannak azoknak a hazai kiberbiztonsági képzéseknek, amelyekhez elsősorban a Nemzeti Közzolgálati Egyetem hallgatói, másodszorban ismeretségi körük tagjai hozzáférnek. A 3. fejezetben már láthattuk, hogy e képzések kiváltképp az aktív védekezések B fődimenziójára, az ismeretek és az eszközök naprakészen tartására készítetik sikeresen a résztvevőket. Ugyanakkor nem volt kimutatható különbség az ilyen képzésen korábban részt vett és a képzésekkel semmilyen kapcsolatba nem került csoport között a tekintetben, mennyit tesznek személyes és privát adataik biztonságáért (A fődimenzió).

A képzések hatásvizsgálatának kérdése azonban nemcsak a konkrét viselkedések, hanem a motiváció szintjén is feltehető. Jellemzően mit váltanak ki e képzések a résztvevőkben az által, hogy ismertetik a legfontosabb visszaéléstípusokat és azok megelőzésének stratégiáit? A képzések hallgatói nagyobb arányban vélik úgy, hogy e fenyegetések valóságok, és őket is érinthetik? Súlyosabbnak látják e visszaélések következményeit, mint korábban? Meggyőzi őket a képzés arról, hogy igenis lehet hatékonyan védekezni e veszélyforrások ellen? Módosítanak-e a képzések a résztvevők énképén: a képzés elvégzésével hajlamosabbak-e úgy látni, hogy ők maguk képesek gondoskodni saját online biztonságukról?

A védelemmotiváció összetevőivel való kapcsolatokat vizsgáló négy *t*-próba közül ez esetben is csak egy mutatott szignifikáns hatást, mégpedig ismét az önhatékonyág dimenziójával ($t_{[213]} = 2,722$; $p = 0,007$). Azonban ebben az esetben is gyenge-közepes mértékű hatásról beszélhetünk (Cohen-féle $d = 0,374$), ahogy azt az 5.4. ábra is szemlélteti.



5.4. ábra: A kiberbiztonsági képzések hatása az önhatékonyságra
(a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés

A védelemmotivációval összefüggést nem mutató demográfiai jellemzők

A védelemmotiváció egyik komponensével sem mutatott szignifikáns korrelációt az életkor; a Pearson-féle korreláció abszolút értéke mind a négy esetben $|r| < 0,14$ (n.s.). Tehát nem látunk arra utaló jelet, hogy a fiatalabb és idősebb generációk másképp értékelnék az online fenyegetéseket és a védekezés lehetőségeit. A lakóhely jellege (település mérete) szintén nem befolyásolta a motiváció egyik elemét sem, tehát a fővárosban, vármegyeszékhelyeken vagy megyei jogú városokban, kisebb városokban, illetve községekben, falvakban élők között nem mutatkozott szignifikáns különbség a vélekedésekben. Ehhez hasonlóan az egyetemi hallgatói státusszal rendelkezőknek sem voltak más hiedelmeik, mint akik jelenleg nem rendelkeznek hallgatói jogviszonnyal.

ÖSSZEGZÉS

A kiberbiztonság és az online fenyegetésekkel szembeni védelem fontossága az elmúlt évtizedekben globális szinten szignifikáns módon megnövekedett. A társadalmi és technológiai változások révén a különböző demográfiai csoportok eltérően reagálnak az ilyen fenyegetésekre, ami alapvetően befolyásolja a védelem hatékonyságát. Az online fenyegetésekkel szembeni attitűdök – például a fenyegetések észlelése, az önhatékonyság érzése és a védekezésre való hajlandóság – kulcsfontosságúak az egyének kiberbiztonsági magatartásának megértésében.

Mindezekkel összefüggésben jelen fejezet a demográfiai csoportok közötti különbségeket vizsgálta az online fenyegetésekkel szembeni védelemmotiváció tükrében. Az eredmények tekintetében célunk volt, hogy feltárjuk a demográfiai különbségek szerepét a védelemmotiváció kialakulásában. Ennek keretében különös figyelmet fordítottunk a nemek, az iskolai végzettség, az életkor és az aktív kereső státusz hatásaira, valamint ezek kapcsolatára az önhatékonyság és más motivációs tényezők alakulásával.

Az eredmények szerint a nemek között eltérés mutatkozott: a nők súlyosabbnak ítélték az online visszaélések következményeit, mint a férfiak. Az iskolai végzettség és az önhatékonyság között szignifikáns kapcsolatot találtunk, a magasabb végzettségűek nagyobb technológiai önhatékonyságot mutattak. Az aktív keresők szintén magabiztosabbak voltak a védekezési képességeikben, mint a nem keresők.

Az életkor, a lakóhely és az egyetemi hallgatói státusz hatásának tekintetében a demográfiai jellemzők nem mutattak szignifikáns kapcsolatot a védelemmotivációval, ami arra utal, hogy ezek az attribútumok nem befolyásolják jelentősen az online fenyegetésekkel szembeni attitűdöket.

A kiberbiztonsági képzések főként az önhatékonyságra gyakoroltak pozitív hatást, míg más motivációs tényezőkre kevésbé. Az életkor, a lakóhely típusa és az egyetemi hallgatói státusz nem mutatott összefüggést a védelemmotivációval. Az eredmények alapján az oktatási programok demográfiai érzékenyítése, különösen a nemek és az iskolai végzettség figyelembevételével, hatékonyan növelheti az információbiztonsági tudatosságot.

A kutatás eredményei alapján a demográfiai tényezők, különösen a nemek és az iskolai végzettség, fontos szerepet játszanak az online védelemmotiváció alakulásában. Az eredmények alátámasztják, hogy az oktatási programok és kiberbiztonsági stratégiák fejlesztésénél figyelembe kell venni e csoportok különböző igényeit. A nemek közötti különbségek például arra utalnak, hogy a nők esetében hangsúlyosabban kell megjeleníteni az egyéni védekezés lehetőségeit, míg a férfiak számára a fenyegetettség komolyságának tudatosítása lehet célravezető. Az iskolai végzettséghez kötött eltérések pedig az oktatási anyagok differenciálását indokolják.

6. A VÉDELEMMOTIVÁCIÓ SZEREPE A VÉDEKEZŐ VISELKEDÉSFORMÁK KIALAKULÁSÁBAN

A korábbi fejezetek részletesen bemutatták a védelemmotiváció elméleti alapjait, kiemelve annak egyes összetevőit és mechanizmusait. Az eddigiekben megismerkedtünk azokkal a pszichológiai tényezőkkel, amelyek hatást gyakorolnak az egyén biztonsági fenyegetettség észlelésére és az arra adott reakciókra. Ebben a fejezetben a védelemmotiváció konkrét szerepét vizsgáljuk a védekező viselkedés kialakításában, különös tekintettel a kiberbiztonsági kontextusra. Célunk, hogy bemutassuk, hogyan ösztönzik a különböző védelemmotivációs tényezők az egyént a megfelelő biztonsági intézkedések alkalmazására, és milyen módon segíti a motiváció a kockázatok elhárítására irányuló attitűdök és cselekvések erősítését. A védelemmotiváció szerepének megértése alapvető fontosságú ahhoz, hogy hatékony stratégiákat dolgozzunk ki az online térben történő tudatos, proaktív védekezés érdekében.

A KIBERBIZTONSÁGI ASPEKTUS

Az emberi hiba – a nem szándékos cselekedetek vagy a cselekvés hiánya – jelentős szerepet játszik a legtöbb kiberbiztonsági és adatvédelmi incidensben. Többek között az emberek biztonságtudatosságának hiánya az emberi hibák egyik figyelemre méltó oka. A biztonságtudatosság magában foglalja a kibertérből érkező kockázatok megértését, amelyek potenciálisan fenyegetést jelenthetnek, incidensekhez vezethetnek, azok megelőzési és enyhítési stratégiáit, és ami még fontosabb, a hozzáállást és a képességet arra, hogy a megszerzett ismereteket megfelelő cselekvéssé vagy viselkedéssé alakítsák a felhasználók. Ebből fakadóan minden biztonságtudatosságot erősítő kezdeményezésnek arra kell törekednie, hogy pozitív változásokat

idézzon elő a felhasználók kiberbiztonsági ismereteiben, hozzáállásában és viselkedésében.⁵⁷ A biztonságtudatos viselkedéssel összefüggésben a védelemmotivációt illetően a sebezhetőség úgy definiálható, mint egy fenyegetés előfordulásának valószínűsége, továbbá, hogy az egyén valószínűleg ki van téve a fenyegetésnek. A fenyegetés súlyosságát a fenyegetés megnyilvánulása miatt bekövetkező negatív következmények súlyosságaként határozzák meg. Ez a fenyegetés megvalósulása esetén fellépő következmények súlyosságát jelenti. Ha az egyén úgy érzékeli, hogy sebezhető a számítógépes biztonságot vagy az okostelefonokat fenyegető veszélyekkel szemben, és hogy az ilyen fenyegetések következményei károsak számára, akkor ez azt eredményezi, hogy egyre tudatosabban ügyel a számítógépe és az okostelefonja biztonságára. A megküzdési értékelés az a kognitív folyamat, amely lehetővé teszi az emberek számára, hogy tudatosabban kezeljék a kiberbiztonsági kérdéseket. Az önhatékonyság az egyén képessége a kiberbiztonsági gyakorlat megvalósítására. A válaszhatékonyság a kiberbiztonsági kontextusban az a felfogás, hogy a biztonságtudatos viselkedés előnyös az egyén számára.⁵⁸

VISELKEDESI FORMÁK A KIBERTÉRBEN

A viselkedés vizsgálata a kiberbiztonsággal kapcsolatban kiemelten fontos, mert az emberek döntései és cselekedetei jelentősen befolyásolják a rendszerek biztonságát. Még a legfejlettebb technológiai védelem is hatástalan lehet, ha a felhasználók nem tartják be az alapvető biztonsági szabályokat, például a jelszavak kezelésében vagy a gyanús linkek elkerülésében. A kiberfenyegetések gyakran a felhasználói hibákat célozzák meg, így a tudatos viselkedés és a megfelelő motivációk kialakítása kulcsfontosságú a kockázatok csökkentésében. A viselkedés vizsgálatával jobban megérthetjük, hogy milyen pszichológiai és társadalmi tényezők befolyásolják az egyének döntéseit a kiberbiztonsági szabályok betartásával kapcsolatban.

⁵⁷ CHAUDHARY 2024.

⁵⁸ KIRAN et al. 2025.

Ez az ismeret lehetőséget nyújt hatékonyabb oktatási programok, szabályozások és kampányok kidolgozására, amelyek elősegítik a biztonságstudatos kultúra kialakulását.

A „COM-B” modell szerint az, hogy egy viselkedés (például a képernyő lezárása ebédelni induláskor) megtörténik-e vagy sem, három, egymással összefüggő tényezőtől függ:

1. a képességtől (Meg tudják-e csinálni? Tudják-e, hogyan kell?);
2. a lehetőségtől (Van-e lehetőségük a cselekvésre?);
3. a motivációtól (Motiváltak-e a képernyő lezárására?).

A beavatkozás típusa a (nem) viselkedés okától függ – így például, ha a felhasználó képes lezárni a képernyőt, lehetősége van rá, de nem motivált, akkor a beavatkozásoknak a motiváció megteremtésén kell alapulnia (például oktatás, tudatosítás, jutalmazás/büntetés). Amennyiben a kezdeti elemzés azt találta, hogy a felhasználó motivált, de nem tudja, hogyan kell lezárni a képernyőt (képesség), akkor a beavatkozásnak a képzésen (képességfejlesztés) kell alapulnia. Ez a példa azonban azt is mutatja, hogy fontos a mögöttes viselkedés valódi okának azonosítása és a megfelelő beavatkozással való reagálás: az áttekintett kvalitatív tanulmányok közül több esetben a munkavállalók tudták, le kellene zárniuk a képernyőket, és azt is, hogyan kell ezt megtenniük, de nem zárták le őket, mert – az ő sajátos munkakörnyezetükben – úgy érezték, hogy ez a munkatársak iránti bizalom hiányát jelentené. Az ENISA egy jelentése több pontján is kiemelte, hogy a szervezet és a személyzet közötti, valamint a személyzet egymás közötti bizalma fontos a kiberbiztonság szempontjából, de itt oktatási beavatkozásra van szükség: a munkavállalóknak meg kell érteniük, hogy a képernyőzárolás kulcsfontosságú biztonsági szokás, és hogy „ez üzleti, nem személyes ügy”. A valós világbeli bizalom és az online bizalom közötti különbség kiváló példa arra, hogy a szervezeteknek hol kellene célba venniük az oktatást és a képességfejlesztést: az olyan támadások, mint az adathalászat és a *social engineering* kihasználják azt a tényt, hogy a legtöbb ember a fizikai világból származó bizalmi modellekre támaszkodik – fontos megérteni, hogy az online világban hol más a bizalom (nem mindig lehetünk biztosak abban, hogy egy e-mail attól

származik, akinek kiadják magukat), és hogyan reagáljunk megfelelően (például konzultáljunk a biztonságiakkal vagy a kollégákkal, ellenőrizzük a személyazonosságukat különböző csatornákon keresztül, udvariasan utasítsuk vissza a kérést).⁵⁹

Azonban fontos kiemelni, hogy a személyiség, a kognitív és viselkedési jellemzők egyéni különbségei is összefüggenek a biztonság tudatos viselkedéssel. Dawson és Thomson szerint a kognitív képességek és a személyiségjegyek egyéni különbségei kulcsszerepet játszhatnak az információs rendszerek biztonságának garantálásában.⁶⁰ Emellett kiemelendő az is, hogy a kiegészítés, a fáradtság és a stressz megnehezíti a biztonságos működést, ezek a problémák már a Covid-19 előtt is léteztek, és a világjárvány idején a kiberbiztonságban az emberi teljesítménnyel kapcsolatos problémák soha nem látott szintet értek el.⁶¹ Sajnos ezek a dilemmák tartósan fennállnak, és továbbra is pusztítást végeznek a kiberbiztonsági erőfeszítésekben, és a problémák orvoslására alig van lehetőség.⁶² A 6.1. táblázat az egyes problémák meghatározását és példáit tartalmazza.

Ezek a problémák negatívan hatnak a kiberbiztonságra a termelés-csökkenés, a megfelelés elmaradása, a megbecsülés hiánya, az alacsony biztonság tudatosság, az iránytalanság és a letargia formájában. Ezek a tünetek ellensúlyozzák a biztonsági tudatossági programok hatékonyságát. Sajnos a kiberbiztonsággal kapcsolatos kedvezőtlen emberi teljesítmény-problémák nagyrészt alul- és kevéssé kutatottak.⁶³ Az emberi teljesítmény problémáit súlyosbítja a szakképzett munkaerő hiánya, a nem megfelelő erőforrások, a rossz irányítás és a nem hatékony prioritások meghatározása. A kiberbiztonsági szakemberek és informatikusok hosszú órákat dolgoznak nagy igénybevétel mellett a kibertámadások, az adatbiztonság megsértése és a zsarolóvírus-támadások megelőzése érdekében.⁶⁴

⁵⁹ DROGKARIS–BOURKA 2018.

⁶⁰ DAWSON–THOMSON 2018.

⁶¹ NOBLES 2021.

⁶² NOBLES 2022.

⁶³ HULL 2017.

⁶⁴ THOMAS 2024.

6.1. táblázat: Az emberi teljesítménnyel kapcsolatos problémák

Tényező	Definíció	Példa
Kiegészítés	A kiegészítést érzelmi fáradtsággént és kimerültségként, deperszonalizációként, kompetenciahiányként és a munkával kapcsolatos távolságtartásként határozzák meg.	1. Túlerhelési kiegészítés 2. Elhanyagolási kiegészítés 3. Alulteljesítési kiegészítés
Kimerültség	A kimerültség lehet letargia vagy vonakodás a biztonsággal és a megfeleléssel való foglalkozástól, mivel túlterheltek vagyunk/a hatékony, biztonság tudatos viselkedés elsajátítását és gyakorlását gyakran gátolja, hogy korlátozottan tudunk racionálisan dönteni.	1. Azonosítási kimerültség 2. Megfelelési kimerültség 3. Riasztási kimerültség 4. Döntési kimerültség 5. Szabályozási kimerültség
Stressz	A stressz fizikai és érzelmi válasz adott körülményekre. Az akut stressz a harc vagy menekülés koncepciója, amelyben a tünetek rövid távúak, míg az epizodikus stressz az ismételt stresszes eseményeknek való kitettségére vonatkozik csökkentett regenerációs időszakokkal. A krónikus stressz a stresszes helyzeteknek való hosszú távú kitettségéből ered, amely akadályozza az egyén képességeit, és kiegészítéshez vezethet.	1. Akut stressz 2. Epizodikus stressz 3. Krónikus stressz 4. Időhiányos stressz 5. Előrelátási stressz 6. Szituációs stressz

Forrás: NOBLES 2022 alapján

A VÉDELEMMOTIVÁCIÓ ÉS A VÉDEKEZŐ VISELKEDÉS KAPCSOLATA

A védelemmotiváció alapvetően befolyásolja a védekező viselkedés kialakulását, mivel az egyének biztonsági fenyegetettségének az érzékelése és az azokra adott válaszreakciók ezen motivációs folyamatokon alapulnak. Amennyiben valaki nagy fenyegetettséget érzékel, ugyanakkor hisz abban, hogy képes hatékonyan védekezni (önhatékony), és hogy a védekező cselekvés eredményes lesz, nagyobb valószínűséggel fog biztonság tudatos viselkedést tanúsítani. A védelemmotiváció tehát összekapcsolja a fenyegetés észlelését a proaktív cselekvéssel, amely elengedhetetlen a kiberbiztonság növeléséhez. A védekező viselkedés kialakulásához gyakran egy elszenvedett támadás vezet a felhasználókat, amelyek lehetnek kifejezetten felhasználó-orientált támadások. Ezek a kibertámadások olyan speciális rendszerekként

definiálhatók, amelyben a támadó a rendszer felhasználóit keresi és veszi célba ahelyett, hogy közvetlenül magát a rendszert támadná. Ezt a sémát számos okból alkalmazzák, például a tűzfalak és behatolásérzékelő rendszerek megkerülésére. Ez a fajta kibertámadás ugyanis a támadók számára könnyebb módot biztosít arra, hogy hozzáférjenek a célrendszerekhez vagy érzékeny eszközökhöz, például adatbázisokhoz, érzékeny fájlokhoz vagy akár az ipari vállalatok irányítási folyamataihoz. Ennek kivitelezéséhez a támadó a támadás felépítéséhez a felhasználóról érzékeny információkat keres és nyer ki, például az e-mail-címét vagy személyes címét, közösségi hálózati profiljait, vagy akár a legközelebbi és legsebezhetőbb barátaira vonatkozó információkat. Miután a felhasználó veszélybe került, a támadó teljes mértékben átvesszi az irányítást a gépe felett, így a rendszerhez való hozzáférést arra használhatja, hogy rosszindulatú szoftvereket terjesszen a vállalati hálózaton belül, és bizalmas adatokat szerezzen meg. Kiemelendő ezek közül is a *spear phishing* támadás, amely az egyszerű adathalásznál speciálisabb típusú adathalász tevékenység, amikor ugyanazt a rosszindulatú e-mailt a lehető legtöbb embernek küldik el. A *spear phishing* gondosan úgy van megtervezve, hogy egyetlen áldozatot célozzon meg. A hackerek időt szánnak arra, hogy mélyreható kutatást végezzenek a célfelhasználókon, és olyan üzeneteket készítsenek, amelyek relevánsabbak és személyesebbek számukra. Például hamisított hivatalos dokumentumokat hozhatnak létre, amelyek személyes adatokat tartalmaznak. A *spear phishing* gyakran jön hitelesnek tűnő hamisított e-mailek formájában, olyan ürüggyekkel, amelyekkel el akarják nyerni a felhasználó bizalmát, és rávenni őt a védelem csökkentésére és a csatolmányok letöltésére, amelyek betöltik a rosszindulatú programot a felhasználó rendszerébe. Az ilyen típusú támadások kevésbé észlelhetők, és nagyon nehéz ellenük védekezni. Emellett még kiemelendők az olyan adathalász *hamis weboldalak*, amelyek pontosan úgy néznek ki, mint az eredetiek. Ezeket az eredeti oldal forráskódjának (HTML/CSS) másolásával hozzák létre, hogy a tapasztalatlan felhasználókat személyes, hitelesítő vagy pénzügyi adataik megadására rávegyék. Ilyenek például a Facebook, Twitter, banki és vállalati oldalaknak kinéző adathalász oldalak. Az úgynevezett *drive-by* letöltéseknél a felhasználót egy csalárd oldalra vezetik, és arra

csábítják, hogy ingyenes programnak (például vírusirtónak, zenének vagy játéknak) álcázott rosszindulatú szoftvereket töltsön le.⁶⁵

Mindannyian találkozhattunk a fentiekben bemutatott támadási módszerekkel a mindennapok során, amelyek ellen különböző megelőző lépéseket tehetünk. Ezek a különböző aktív védekezési vagy magatartási formák segíthetnek a kibertérből érkező fenyegetésekkel szembeni védekezésben. Bizonyos tanulmányok szerint a biztonságtudatos magatartás több fontos biztonsági viselkedés összességéként mérhető. Ezek az *eszközbiztonsági és frissítő*, a *proaktív* és a *jelszógeneráló magatartás*. Az *eszközbiztonsági és frissítési magatartás* arra a viselkedésre utal, hogy az eszközök lezárásához PIN-kódot és jelszót használnak, automatikus eszköz- vagy képernyőzárat állítanak be, biztonsági javításokat alkalmaznak, vagy más módon naprakészen tartják a szoftverüket, és manuálisan lezárják az eszközöket, mielőtt felügyelet nélkül hagyják őket. A *proaktív magatartás* a kontextuális jelzésekre, például az URL-sávra vagy más böngészőjelzőkre való odafigyelésre utal a weboldalakon vagy e-mail-üzenetekben, az óvatosságra a weboldalakon történő adatszolgáltatás során, valamint a biztonsági incidensek bejelentése során tanúsított proaktív viselkedésre. Végül a *jelszógeneráló magatartás* arra a viselkedésre utal, hogy erős jelszavakat választunk, és nem használjuk újra a jelszavakat különböző fiókok között.⁶⁶

Mindemellett fontos, hogy naprakészek legyünk a minket körülvevő fenyegetésekkel, trendekkel kapcsolatosan. Ennek első lépése, hogy megértsük, miért is fontos a kiberbiztonság. Ehhez fontos kiemelni, hogy a kibertámadások fenyegetettségi köre folyamatosan fejlődik, a kiberbűnözők egyre szervezettebbé válnak, és fejlett technikákat alkalmaznak a biztonsági rendszerek feltörésére. Olyan taktikákat alkalmaznak, mint a korábban már említett adathalász e-mailek, rosszindulatú szoftverek, zsarolóprogramok és a *social engineering*, hogy a gyanútlan embereket érzékeny információk kiadására vagy a rendszerekhez való jogosulatlan hozzáférésre bírják. A kiberbiztonság emellett létfontosságú szerepet játszik a személyes és pénzügyi információk védelmében, hiszen olyan intézkedések és gyakorlatok egész sorát foglalja magában, amelyek segítenek megvédeni

⁶⁵ HAMOUD-ÂÎMEUR 2020.

⁶⁶ HUMAIDI – ABDALLAH ALGHAZO 2022.

az adatokat a jogosulatlan hozzáféréstől, a visszaélésektől és a károktól. Az erős jelszavak használatától kezdve a titkosítás és a tűzfalak alkalmazásáig az az intézkedések célja, hogy naprakészek maradjunk. Ennek megvalósításához elengedhetetlen, hogy folyamatosan tájékozódjunk a legújabb fenyegetésekről, technológiákról és legjobb gyakorlatokról. Szerencsére számos stratégiát alkalmazhatunk, amelyekkel biztosíthatjuk, hogy jól tájékozottak legyünk. Az olyan nemzeti kiberbiztonsági hírforrások, mint például a Nemzetbiztonsági Szakszolgálat, a Nemzeti Kibervédelmi Intézet (NBSZ NKI) megbízható és naprakész információkat nyújt a kiberfenyegetésekről és a legjobb gyakorlatokról.⁶⁷ Weboldaluk és közösségimédia-fiókjaik követése segíthet naprakésznek maradni a kiberbiztonsággal kapcsolatos trendeket illetően. Továbbá számos szervezet és oktatási intézmény kínál webináriumokat és online tanfolyamokat különböző témákban, amelyek lehetővé teszik, hogy elmélyítsük ismereteinket, és tanuljunk az iparág szakértőitől. Amennyiben aktívan részt veszünk ezeken a webináriumokon és programokon, értékes ismeretekre tehetünk szert amellet, hogy naprakész információkat kapunk a kiberbiztonságról.⁶⁸

EREDMÉNYEK

Kutatásunk egyik sarokköve volt az online felületeken tanúsított aktív védekező viselkedési formák motivációs hátterének feltárása. Az ezekre való hajlandóság nagyfokú variabilitást mutat nemcsak abban a tekintetben, ki mennyit tesz a saját biztonságáért, hanem abban is, ki milyen módon védekezik. A védekezésben mutatott egyéni különbségek magyarázata azért fontos feladat, mert ha pontosabban megértjük e viselkedések megjelenésének, illetve hiányának okait, akkor azt is látni fogjuk, mit érdemes tenni annak érdekében, hogy az online biztonságra való törekvés minél inkább elterjedjen az egyetemi hallgatóság, illetve a teljes népesség körében.

⁶⁷ A Nemzeti Kibervédelmi Intézet weboldala (<https://nki.gov.hu/>).

⁶⁸ Institute of Data 2023.

Kutatási célkitűzésünket a fenti alkalmazási lehetőségeken túlmenően egy tisztán elméleti kérdésfelvetés is vezérelte. A védelemmotivációs elmélet az elmúlt négy évtizedben sikeresen magyarázta az emberi védekező viselkedések széles körét. Felvetődik ugyanakkor a kérdés, hogy megfelelően általános érvényű-e az elmélet ahhoz, hogy az online térben és a digitális eszközök használatakor alkalmazott védekező viselkedésekre is jól működő magyarázattal szolgáljon. Az online fenyegetések természete jóval absztraktabb, mint a fizikai világban ránk leselkedő veszélyek. Az az elmélet, amely képes megmagyarázni, miért, vagy éppen miért nem kapcsolják be az emberek a biztonsági öveiket, zárják be a bejárati ajtót vagy kerülnek el bizonyos városrészeket, képes-e egyben arra is, hogy megokolja, miért, vagy miért nem használnak az emberek biztonságos jelszavakat, állítják privátra közösségimédia-profiljukat, vagy indítanak el ellenőrizetlen forrásból származó futtatható állományokat számítógépükön?

Hogy e kérdéskört teljeskörűen feltárjuk, korrelációs és regressziós elemzéseket végeztünk az aktív védekezés különböző szintű mutatói (összesített mutató, fődimenziók, aldimenziók) és a védelemmotiváció négy összetevője között.

Az elemzés első fázisában az úgynevezett nulladrendű korrelációk azt mutatták, hogy az aktív védekezés összesített mutatója a négy motivációs tényező közül kettővel mutat szignifikáns összefüggést, mégpedig azokkal, amelyek a megküzdési stratégiák értékeléséhez tartoznak. A válaszhatékonysággal a korreláció mértéke $r = 0,217$ ($p = 0,001$), míg az önhatékonyság esetében jóval erősebb: $r = 0,368$ ($p < 0,001$). Amikor azonban egy többváltozós regressziós modellben az aktív védekezést e két tényezőtől együttesen próbáljuk megjósolni, az önhatékonyság mellett a válaszhatékonyságnak nincsen további magyarázó ereje, az önhatékonyság marad az egyetlen szignifikáns tényező. E mintázat úgy értelmezhető, hogy a válaszhatékonyság hatása az aktív védekező viselkedési formákra nem közvetlenül, hanem az önhatékonyságon mint mediátorváltozón keresztül érvényesül.⁶⁹ Ez az eredmény azzal a lehetséges oksági modellel is összhangban áll, amelyet

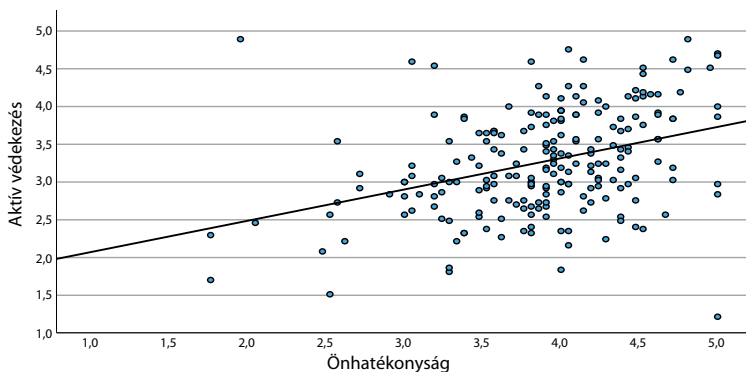
⁶⁹ BARON-KENNY 1986.

a védelemmotiváció összetevőiről szóló 4. fejezet végén ismertettünk: e szerint az oksági hatások láncolata a válaszhatékonysággal kezdődhet, ez hat pozitívan az önhatékonyságra, amely ezt követően csökkenti a személyes veszélyeztettség érzetét, majd – az utolsó lépésben – ennek hatására csökken a fenyegetés észlelt súlyossága is. Az itt bemutatott eredmény a modellünket kiegészíti azzal, hogy e láncolaton belül az önhatékonyság kulcsfontosságú tényező, hiszen ez az a pont, ahonnan az online platformokon mutatott aktív védekező viselkedések erednek.

Ennek az eredménynek fontos alkalmazott vonatkozásai vannak, hiszen egyértelműen azt mutatja, hogy például egy kiberbiztonsági képzés megtervezésekor nem feltétlenül kell arra törekedni, hogy a résztvevők belássák, milyen gyakoriak ezek a visszaélések, és milyen súlyos következményekkel járhatnak. A személyes veszélyeztettség és a fenyegetés észlelt súlyossága növekedésének természetesen lesz kimutatható hatása – például a résztvevő aggodalmaskodni (rosszabb esetben szorongani) kezd. Eredményeink szerint azonban semmilyen mértékű aggodalom vagy szorongás nem képes kiváltani azt, hogy az egyén valóban elkezdje megtenni a szükséges lépéseket online biztonságáért. Ehhez a változáshoz egyetlen út vezet, ha beláttatjuk a résztvevővel, hogy ő maga igenis képes hatékonyan megvédeni magát a fenyegetések ellen.

Ezen eredmény tükrében különösen érdekes az 5. fejezetben bemutatott részeredmény, amely szerint a kiberbiztonsági képzések hatására éppen az önhatékonyság emelkedik, míg e képzések a védelemmotiváció egyéb komponenseire nincsenek detektálható hatással. Vagyis adataink azt mutatják, hogy a jelenleg működő és az egyetemi hallgatók számára hozzáférhető kiberbiztonsági képzéseink éppen azt a típusú hatást fejtik ki, amelyre szükség van. Ahelyett, hogy tévesen riogatásra vagy félelemkeltésre apellálnának, és ahelyett, hogy misztifikálnák a védekezés lehetőségeit, konkrét útmutatást nyújtva meggyőzik a résztvevőket arról, hogy saját maguk képesek megtenni mindent, ami a biztonságuk érdekében szükséges.

Az önhatékonyság és az aktív védekezés összesített mutatója közötti korrelációt a 6.1. ábrán látható pontdiagram szemlélteti:



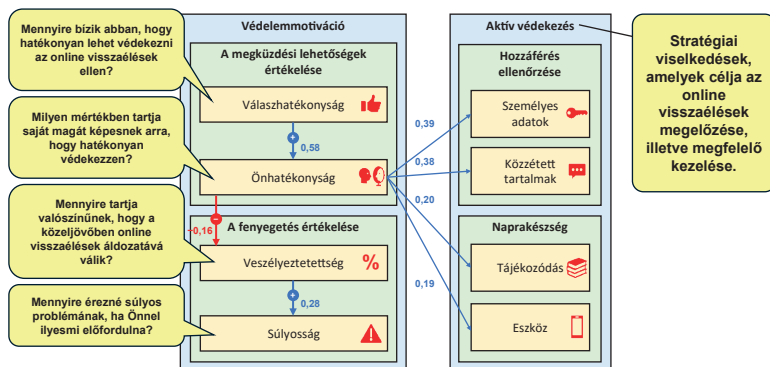
6.1. ábra: Az önhatékonyság és az aktív védekező viselkedési formák összesített mutatója közötti korreláció

Forrás: saját szerkesztés

Az ábrán látható regressziós egyenes egyenlete: $\text{aktív védekezés} = 1,66 + 0,414 \times \text{önhatékonyság}$, vagyis a mért aktív védekezés várható értéke 0,414 ponttal emelkedik, ha az önhatékonyság 1 ponttal nő (mindkét változót 1-től 5-ig terjedő skálán mérjük).

Az eredményeket tovább erősíti, hogy pontosan ugyanezt a mintázatot kapjuk, ha az elemzést megismétljük az aktív viselkedés fő-, illetve aldimenzióinak szintjén is. Az önhatékonyság minden esetben szignifikánsan meghatározza az adott viselkedéstípus gyakoriságát, és ha többváltozós regressziós modellt illesztünk az adatsorra, akkor minden esetben az önhatékonyság marad az egyetlen szignifikáns prediktor. (Önmagában véve a válaszhatékonyság is szignifikánsan előre jelzi az aktív védekezés A fődimenzióját és annak mindkét aldimenzióját, de ez a kapcsolat azonnal eltűnik, ha a válaszhatékonyság az önhatékonyság mellett szerepel a modellben. Vagyis a válaszhatékonyság hatása ezekben az esetekben is indirekt módon, az önhatékonyságra gyakorolt hatásán keresztül érvényesül.)

Az önhatékonyság egyes aldimenziókra gyakorolt hatását a 4. fejezetben vázolt kauzális védelemmotiváció-moddellel együtt jeleníti meg a 6.2. ábra. A nyilak mellett szereplő számok minden esetben egyszerű (nulladrendű) korrelációk.



6.2. ábra: A védelemmotiváció és az aktív védekezés aldimenziói között kimutatott összefüggésrendszer (a feltüntetett értékek Pearson-féle korrelációs együtthatók)

Forrás: saját szerkesztés

A fentiekén túlmenően az ábráról az is leolvasható, hogy az önhatékonyság szerepe jóval nagyobb a személyes adatok és privát tartalmak védelméhez és ellenőrzéséhez kapcsolódó viselkedések megjelenésében (A fődimenzió), mint a rendszeres tájékozódást és az eszközök naprakészen tartását szolgáló viselkedések kialakulásában (B fődimenzió), bár a hatás ez utóbbi esetben is szignifikáns. Ugyanakkor a 3. fejezetben láttuk, hogy a kiberbiztonsági képzések kimutathatóan megnövelik a B fődimenzióhoz tartozó viselkedések gyakoriságát. Az ebben a fejezetben bemutatott eredmények ezen túlmenően fényt derítettek arra a részletre, hogy e korábban tárgyalt hatást a képzések nem a védelemmotiváció kiváltásán keresztül érik el (nem azáltal, hogy kialakítják a résztvevőben a meggyőződést, hogy ő „igenis képes tájékozódni és rendszeresen frissíteni az eszközeit”), hanem közvetlen módon (vagy esetleg más mediátorváltozón keresztül, amely nem képezi jelen vizsgálatunk tárgyát).

ÖSSZEGZÉS

A fejezetben a védelemmotiváció és a kiberbiztonsági védekező viselkedés közötti kapcsolat vizsgálatát mutattuk be, különös tekintettel az egyéni viselkedésformákra és az azok háttérében meghúzódó pszichológiai mechanizmusokra. A bevezetésben megfogalmazott alapvető hipotézis szerint a védelemmotiváció, amely a fenyegetés súlyosságának, a sebezhetőségnek, valamint az ön- és válaszhatékonyságnak az észlelésére épül, alapvetően meghatározza az egyén védekező magatartását az online térben.

Az eredmények alátámasztották, hogy az önhatékonyság és a válaszhatékonyság a legjelentősebb előrejelzői az aktív védekező viselkedésnek. Az önhatékonyság hatása különösen erős volt, ami azt mutatja, hogy az egyének motivációja a saját képességeikbe vetett hitükön alapul. A válaszhatékonyság ugyan közvetlenül nem bizonyult szignifikánsnak többváltozós modellekben, de közvetetten, az önhatékonyságra gyakorolt hatásán keresztül befolyásolta a védekező viselkedést. A bevezetésben hangsúlyozott másik két tényező – a fenyegetés súlyossága és a sebezhetőség érzékelése – kisebb szerepet játszott az aktív védekezési stratégiák kialakításában. Ezek inkább az aggodalom és a kockázatészlelés szintjét növelték, de nem eredményeztek közvetlen változásokat a viselkedésben. Ez arra utal, hogy a félelemkeltés helyett a képességek fejlesztésére kell helyezni a hangsúlyt a kiberbiztonsági oktatások során.

A kutatás az aktív védekezési stratégiákat négy dimenzió mentén vizsgálta:

- A1: Személyes adatok.
- A2: Közzétett tartalmak.
- B1: Tájékozódás.
- B2: Eszközök.

Az eredmények szerint a személyes adatok védelmére irányuló viselkedések a leggyakoribbak. Például a résztvevők többsége erős jelszavakat használ, és kerüli a személyes adatok nyilvános megosztását. Az információmegosztás ellenőrzése szintén magas szintet mutatott, különösen a közösségi média

beállításában és a privát kommunikációs csatornák használatában. Ezzel szemben a naprakészséggel és frissítésekkel kapcsolatos viselkedések kevésbé gyakoriak. Például a résztvevők csupán 30%-a olvassa el rendszeresen az adatvédelmi irányelveket, és kevesen követik aktívan a kiberbiztonsági híreket. Ez alátámasztja a bevezetésben megfogalmazott aggodalmat, hogy a tájékozottság hiánya sebezhetőbbé teszi az egyéneket. Az önhatékonyság szignifikáns pozitív hatást gyakorolt mind az A (adatvédelem) fődimenzióra, mind a B (naprakészség és frissítések) fődimenzióra. Ez azt jelzi, hogy az egyének nagyobb valószínűséggel tesznek lépéseket a biztonságuk érdekében, ha úgy érzik, hogy képesek hatékonyan cselekedni. Ugyanakkor a B dimenzióban a hatás gyengébb volt, ami viszont azt mutatja, hogy ezekhez a viselkedésekhez olyan további motivációs tényezők szükségesek, mint például az oktatás vagy a közösségi normák erősítése.

A bevezetésben kiemelt kiberbiztonsági oktatás és tudatosság növelésének szükségességét az eredmények teljes mértékben alátámasztják. A bevezetésben felvázolt problémák – például az információhiány és az emberi hiba – az eredményekben is megjelennek, különösen a naprakészség és a rendszeres frissítések alacsony szintjén keresztül. Az eredmények ugyanakkor rámutatnak, hogy a személyes adatok védelme már jobban beépült a felhasználói szokásokba.

Az eredmények alapján az alábbiakat javasolhatjuk:

- Oktatás: az önhatékonyság fejlesztésére irányuló tréningek és kampányok megalkotását szorgalmazzuk, amelyek gyakorlati útmutatást adnak a kiberbiztonsági intézkedések megvalósításához.
- Kommunikáció: a kiberbiztonsági kockázatokról szóló tájékoztatás növelését tanácsoljuk, különös tekintettel a naprakész információk és hírek terjesztésére.
- Technológiai megoldások: könnyen használható jelszókezelők és frissítési rendszerek népszerűsítését indítványozzuk, amelyek csökkentik a felhasználók terheit.

A kutatás egyértelművé tette, hogy a védelemmotiváció, különösen az önhatékonyság, kulcsszerepet játszik az aktív védekezési stratégiák kialakításában. Míg a személyes adatok védelme terén jelentős előrelépés történt, a tájékozottság és a frissítések terén további intézkedések szükségesek. Az oktatási programoknak a gyakorlati képességek fejlesztésére kell összpontosítaniuk, miközben elősegítik a kiberbiztonsági kultúra elmélyítését.

7. A SZEMÉLYISÉG ÉS VÉDELEMMOTIVÁCIÓ DIMENZIÓINAK ÖSSZEFÜGGÉSRENDSZERE

A SZEMÉLYISÉG ÉS A VISELKEDÉS ÖSSZEFÜGGÉSEI

A személyiség olyan stabil pszichológiai vonások rendszere, amely befolyásolja az egyének érzékelését, érzelmeit és viselkedését különböző helyzetekben.⁷⁰ A Nagy Ötök (*Big Five*) személyiségmodell öt alapvető dimenzióban ragadja meg az emberi személyiség főbb vonásait: nyitottság az élményekre, lelkiismeretesség, extravertzió, barátságosság és neuroticizmus formájában. Ez a modell széles körben alkalmazható az emberi viselkedés megértésében, beleértve az online környezetben tapasztalható magatartásformák és védekezési stratégiák elemzését is.

Az online térben tapasztalható védekező viselkedések – mint például a jelszókezelés, a kétfaktoros hitelesítés alkalmazása vagy a közösségi média adatvédelmi beállításainak rendszeres ellenőrzése – nem kizárólag a személyiségjegyek függvényei. A védekezési döntések sokkal inkább a fenyegetés észlelésén, a védekezés hatékonyságába vetett hiten és az egyének motivációs rendszerén alapulnak. Az egyik alapvető elmélet, amely ezt a folyamatot magyarázza, a védelemmotiváció elmélet (továbbiakban: *protection motivation theory*, PMT).⁷¹ Ez az elmélet arra utal, hogy az emberek akkor hajlamosak védekező intézkedéseket tenni, ha a fenyegetés mértékét súlyosnak ítélik meg, és hatékonynak tartják azokat a stratégiákat, amelyekkel a veszélyeket elkerülhetik.

Elemzésünkben abból indulunk ki, hogy a személyiségjegyek nem közvetlenül befolyásolják az online védekező viselkedést, hanem a védelemmotiváció különböző összetevőinek (fenyegetettségészlelés, a védekezés hatékonyságába vetett hit és a cselekvési hajlandóság) kialakulására

⁷⁰ MCCRAE–COSTA 1997.

⁷¹ ROGERS 1975.

és felerősödésére hajlamosítanak. Más szavakkal, a személyiségjegyek olyan keretet biztosítanak, amely meghatározza, hogy az egyének milyen mértékben észlelik a fenyegetéseket, hogyan értékelik a védekezési lehetőségeiket, és milyen mértékben motiváltak ezek alkalmazására.

A NAGY ÖTÖK MODELL DIMENZIÓI ÉS A VÉDELEMMOTIVÁCIÓ ÖSSZETEVŐI

A védelemmotiváció elmélete szerint az egyének védekező viselkedését három fő tényező befolyásolja:

1. Fenyegetettségészlelés – annak megítélése, hogy a veszély milyen mértékben jelent kockázatot.
2. Védekezési hatékonyságba vetett hit – annak felismerése, hogy a rendelkezésre álló stratégiák mennyire képesek csökkenteni a kockázatot.
3. Cselekvési hajlandóság – az egyén tényleges motivációja arra, hogy tudatosan védekezzen.

Az alábbiakban részletezzük, hogyan befolyásolják a Nagy Ötök modell dimenziói ezeket az összetevőket.

1. Nyitottság az élményekre (*openness to experience*)

A nyitottság az új tapasztalatokra és kreatív gondolkodásra való hajlamot tükrözi. Azok, akik magas pontszámot érnek el ebben a dimenzióban, nyitottak az innovációkra, és gyakran próbálnak ki új technológiákat vagy online platformokat.

- ♦ Fenyegetettségészlelés: a nagyon nyitott egyének hajlamosak alábecsülni az online fenyegetések súlyosságát, mivel kíváncsiságuk és az újdonság iránti vágyuk elnyomhatja a kockázatok érzékelését.
- ♦ Védekezési hatékonyság: ők gyakrabban próbálnak ki új védelmi eszközöket, például innovatív jelszókezelőket vagy VPN-szolgáltatásokat, ami növeli védekezési eszköztárukat.

- Cselekvési hajlandóság: a kockázatészlelés hiánya miatt alacsonyabb lehet a cselekvési hajlandóságuk a meglévő védekezési protollok alkalmazásában.

2. Lelkiismeretesség (*conscientiousness*)

A lelkiismeretesség a szervezett, felelősségteljes és célorientált magatartást tükrözi. Az ilyen személyek általában alaposak és elővigyázatosak a mindennapi életükben, ami kiterjed az online környezetre is.

- Fenygetettségészlelés: a lelkiismeretes egyének nagyobb valószínűséggel érzékelik a fenyegetéseket, mivel hajlamosak az alapos információgyűjtésre és a potenciális veszélyek részletes elemzésére.
- Védekezési hatékonyság: az erős célorientáltság miatt ezek az egyének bíznak abban, hogy az általuk alkalmazott stratégiák hatékonyak, és következetesen alkalmazzák azokat.
- Cselekvési hajlandóság: a lelkiismeretesség szorosan összefügg az olyan proaktív védekező viselkedéssel, mint például a jelszavak gyakori megváltoztatása vagy az adatvédelmi beállítások rendszeres ellenőrzése.

3. Extraverzió (*extraversion*)

Az extrovertált egyének szociális aktivitásuk és energikus természetük miatt gyakran intenzíven használják a közösségi médiát és más online platformokat, ami növeli kockázati kitettségüket.

- Fenygetettségészlelés: az extrovertált egyének általában kevésbé érzékelik a fenyegetéseket, mivel inkább a társasági élményeket helyezik előtérbe, mint a biztonságot.
- Védekezési hatékonyság: bár képesek lehetnek az egyszerű védekezési stratégiák elfogadására, például az alapvető jelszókezelési szabályok betartására, hajlamosak alábecsülni a bonyolultabb védekezési eszközök szükségességét.

- ♦ Cselekvési hajlandóság: a kockázatkereső beállítottságuk miatt kevésbé hajlamosak elköteleződni a következetes védekező magatartások mellett.

4. Barátságosság (*agreeableness*)

A barátságosság dimenziójába tartozó személyek jól tudnak együttműködni, és könnyebben megbíznak másokban, ami az online térben fokozott sebezhetőséghez vezethet.

- ♦ Fenygetettségészlelés: a barátságos egyének gyakran alábecsülik az online fenyegetéseket, mivel hajlamosak jóindulatúan értelmezni mások szándékait.
- ♦ Védekezési hatékonyság: az együttműködési hajlamuk miatt könnyebben elfogadják a biztonsági tanácsokat, bár nem feltétlenül értik meg azok jelentőségét.
- ♦ Cselekvési hajlandóság: az alacsony fenygetettségészlelés miatt gyakran elmarad a következetes védekezési magatartás.

5. Neuroticizmus (*neuroticism*)

A neuroticizmus dimenziója az érzelmi stabilitást vizsgálja; az érzelmileg stabil egyének általában higgadtan reagálnak a fenyegetésekre, míg az érzelmileg instabil (magas neuroticizmusszinttel rendelkező) személyek hajlamosak túlreagálni azokat.

- ♦ Fenygetettségészlelés: az érzelmileg stabil egyének objektív módon ítélik meg a veszélyeket, míg a neurotikus egyének hajlamosak a fenyegetések eltúlzására.
- ♦ Védekezési hatékonyság: az érzelmileg stabil egyének bíznak a védekezési stratégiák sikerességében, míg a neurotikusok esetében a magas stressz-szint csökkentheti a védekezési intézkedések elfogadását.
- ♦ Cselekvési hajlandóság: az érzelmi instabilitás miatt elkerülő viselkedés vagy impulzív döntések akadályozhatják a hatékony védekezést.

Összességében a Nagy Ötök személyiségdimenziói jelentős hatással vannak a védelemmotiváció különböző összetevőinek alakulására, és ezáltal közvetetten befolyásolják az online védekező viselkedéseket. Az eredmények alapján a személyiség figyelembevétele kulcsfontosságú a célzott kiberbiztonsági stratégiák kidolgozásában. Az ilyen stratégiák hatékonyan támogathatják az egyéneket abban, hogy személyiségüknek megfelelő védekezési módszereket alkalmazzanak, így növelve a digitális tér biztonságát.

MÓDSZER

Az öt, átfogó személyiségdimenziót Maples-Keller és szerzőtársai validált kérdőíve, az IPIP-NEO-60 magyar fordításával mértük.⁷² A tanulmány a kérdőív összeállításának módszertanát részletesen ismerteti, igen alapos, sokrétű pszichometriai elemzéssel demonstrálja az eszköz megbízhatóságát és érvényességét. A IPIP-NEO-60 mindössze hatvan tételből áll, tehát jóval tömörebb a korábbi, jellemzően 300 tételből álló klasszikus személyiségkérdőíveknél; ezért is jelentős módszertani eredmény, hogy a szerzők olyan mérőeszközt alkottak, amely jóval kevesebb információból képes nagy közelítéssel azonos eredményeket adni, mint a jóval terjedelmesebb – és így messze idő- és erőforrás-igényesebb – eszközök. Az IPIP-NEO-60-ban minden személyiségdimenziót 12-12 tétel mér, amelyeken belül 2-2 tétel vonatkozik az egyes személyiségdimenziókat alkotó 6-6 fazettára.

A kutatásban használt személyiségblokk közvetlenül megfelel Maples-Keller és szerzőtársai kérdéssorának, egyetlen item kivételével, amely az eredeti kérdőívben a liberalizmus („O6”) fazettát méri az „I tend to vote for liberal candidates” („Liberális jelöltekre szoktam szavazni”) állítással. Úgy éreztük, hogy kulturális és politikai okok miatt e mondat fordítása a magyar adatközlők számára más konnotációjú lehet, mint a kérdőív

⁷² MAPLES-KELLER et al. 2019.

validálásában közreműködő amerikai adatközlők esetében, ezért potenciálisan veszélyeztetheti a mérés érvényességét. Ezért ezt a tételt egy másik (az eredetivel ellentétben fordított) tétellel helyettesítettük, amely ugyancsak a „liberalizmus” fazettát méri, Johnson szintén IPIP-alapú kérdőívében szerepel, és hasonlóan jó pszichometriai jellemzői vannak: „I believe we should be tough on crime” („Úgy gondolom, hogy a bűnözést keményen kell büntetni”).⁷³

A személyiségjegyekre vonatkozó szakaszt a kérdőívünkben a következő instrukció vezette be:

Kutatásunkban azt is vizsgáljuk, hogyan függnek össze a fentiekben megfogalmazott vélekedések, aggodalmak és szokások a felhasználó személyiségének különféle jegyeivel. A következő blokkban azt kérjük, nyilatkozzon arról, a saját énképe szerint milyen mértékben jellemzők Önre az alábbi állítások.

- 1 – Egyáltalán nem jellemző rám.
- 2 – Inkább nem jellemző rám.
- 3 – Változó: nem tudom eldönteni.
- 4 – Inkább jellemző rám.
- 5 – Nagyon jellemző rám.

Az 7.1. táblázat az öt fő domént, a fazettákat és az azokhoz kapcsolódó tételek magyar fordításait tartalmazza, amelyeket saját kérdőívünkben alkalmaztunk. Az előjel azt jelzi, hogy a tétel megfogalmazása egyenes („+”) vagy fordított-e („-”), azaz hogy az egyetértés az adott domén/fazetta esetében a címke által kifejezett tulajdonság (például „lelkiismeretesség”) vagy annak hiányát („lelkiismeretlenség”) fejezi ki.

⁷³ JOHNSON 2014.

7.1. táblázat: Az öt fő személyiségdimenziót mérő tételek

Domén	Fazetta	Tétel szövege	Előjel	Pozíció
neuroticizmus	N1: szorongás	Aggodalmaskodó vagyok.	+	35
neuroticizmus	N1: szorongás	Könnyen ideges leszek.	+	1
neuroticizmus	N2: ingerlékenység	Könnyen dühbe gurulok.	+	6
neuroticizmus	N2: ingerlékenység	El szoktam veszteni a türelmemet.	+	40
neuroticizmus	N3: depresszió	Gyakran vagyok lehangolt.	+	14
neuroticizmus	N3: depresszió	Nem szeretem magam.	+	45
neuroticizmus	N4: féltékenység	Könnyen meg tudok ijedni.	+	16
neuroticizmus	N4: féltékenység	Nehezen közeledem másokhoz.	+	48
neuroticizmus	N5: mértéktelenség	Ritkán esem túlzásokba.	–	21
neuroticizmus	N5: mértéktelenség	Képes vagyok úrrá lenni a vágyaimon.	–	54
neuroticizmus	N6: sebezhetőség	Nyomás alatt is higgadt maradok.	–	28
neuroticizmus	N6: sebezhetőség	Feszült helyzetekben is nyugodt vagyok.	–	58
extraverzió	E1: kapcsolatteremtő készség	Könnyen barátkozom.	+	32
extraverzió	E1: kapcsolatteremtő készség	Társaságban nyugodtan és magabiztosan viselkedem.	+	5
extraverzió	E2: társaságkedvelés	Szeretem a nagy összejöveteleket.	+	37
extraverzió	E2: társaságkedvelés	Kerülöm a tömeget.	–	9
extraverzió	E3: asszertivitás	Át szoktam venni az irányítást.	+	12
extraverzió	E3: asszertivitás	Próbálok másokat irányítani.	+	42
extraverzió	E4: aktivitás	Mindig elfoglalt vagyok.	+	19
extraverzió	E4: aktivitás	Állandóan rohanok.	+	46
extraverzió	E5: kalandvágy	Szeretem az izgalmakat.	+	53
extraverzió	E5: kalandvágy	Keresem a kalandot.	+	25
extraverzió	E6: vidámság	Sok mindennek tudok örülni.	+	56
extraverzió	E6: vidámság	Szeretem az életet.	+	29

Domén	Fazetta	Tétel szövege	Előjel	Pozíció
nyitottság	O1: képzelőerő	Élénk képzelőerővel rendelkezem.	+	31
nyitottság	O1: képzelőerő	Szeretek álmodozni.	+	4
nyitottság	O2: művészi érdeklődés	Hiszek a művészet fontosságában.	+	36
nyitottság	O2: művészi érdeklődés	Nem kedvelem a művészetet.	–	10
nyitottság	O3: emocionalitás	Intenzíven élem meg az érzelmeimet.	+	43
nyitottság	O3: emocionalitás	Nem nagyon befolyásolnak az érzelmeim.	–	11
nyitottság	O4: érdeklődés	Inkább ragaszkodom az olyan dolgokhoz, amiket már ismerek.	–	47
nyitottság	O4: érdeklődés	Nem szeretem a változás gondolatát.	–	18
nyitottság	O5: intellektus	Kerülöm a magasröptű beszélgetéseket.	–	22
nyitottság	O5: intellektus	Nem érdekelnek az elméleti viták.	–	55
nyitottság	O6: liberalizmus	Az egyetlen igaz vallásban hiszek.	–	60
nyitottság	O6: liberalizmus	Úgy gondolom, hogy a bűnözést keményen kell büntetni.	–	27
barátságosság	A1: bizalom	Megbízom másokban.	+	2
barátságosság	A1: bizalom	Hiszem, hogy az emberek jó szándékúak.	+	34
barátságosság	A2: erkölcsösség	Nem riadok vissza a csalástól, hogy előbbre jussak.	–	38
barátságosság	A2: erkölcsösség	Kihasználok másokat.	–	7
barátságosság	A3: altruizmus	Szeretek másoknak segíteni.	+	41
barátságosság	A3: altruizmus	Szívemen viselem mások dolgait.	+	13
barátságosság	A4: együttműködés	Megsérték embereket.	–	49
barátságosság	A4: együttműködés	Megbosszulom, ha valaki ártott nekem.	–	20
barátságosság	A5: szerénység	Úgy gondolom, hogy jobb vagyok másoknál.	–	23
barátságosság	A5: szerénység	Sokra tartom magam.	–	51
barátságosság	A6: együttérzés	Együttérzek a hajléktalanokkal.	+	57

Domén	Fazetta	Tétel szövege	Előjel	Pozíció
barátságosság	A6: együttérzés	Együttérzek azokkal, akik nálam rosszabb helyzetben vannak.	+	26
lelkiismeretesség	C1: önhatékonyság	Ügyesen intézem a dolgokat.	+	3
lelkiismeretesség	C1: önhatékonyság	Tudom, hogyan kell a dolgokat megoldani.	+	33
lelkiismeretesség	C2: rendszeretet	Szeretem a rendet.	+	8
lelkiismeretesség	C2: rendszeretet	Rendetlenséget szoktam hagyni a szobámban.	–	39
lelkiismeretesség	C3: kötelességtudat	Igazat szoktam mondani.	+	15
lelkiismeretesség	C3: kötelességtudat	Nem tartom be az ígéreteimet.	–	44
lelkiismeretesség	C4: szorgalom	Keményen dolgozom.	+	50
lelkiismeretesség	C4: szorgalom	Sokat várok el magamtól és másoktól.	+	17
lelkiismeretesség	C5: önfegyelem	Megvalósítom a terveimet.	+	52
lelkiismeretesség	C5: önfegyelem	Nehezen fogok hozzá a feladatokhoz.	–	24
lelkiismeretesség	C6: óvatosság	Átgondolatlan döntéseket hozok.	–	30
lelkiismeretesség	C6: óvatosság	Meggondolatlanul cselekszem.	–	59

Forrás: MAPLES-KELLER et al. 2019 alapján

A 7.1. táblázatban az áttekinthetőség érdekében a tételeket domének és fazetták szerint rendezve tüntettük fel. A pozíció oszlopban látható, hogy a kérdőívblokkban az adott tétel hányadik helyen szerepelt. A kérdőív tételek sorrendjének meghatározásakor az alábbi lépéseket követtük: minden fazetta esetében a tételek egyikét véletlenszerűen a sorozat első feléhez (1–30. pozíció), a másik tételt pedig a sorozat második feléhez (31–60. pozíció) rendeltük. Ezt követően mindkét 30-as sorozaton belül a tételeket fazettaindex (a domén betűjelét követő szám) szerint, majd ezen ötelemű csoportokon belül véletlenszerűen sorba rendeztük. Az így előálló sorrend biztosítja, hogy minden egyes fazettára vonatkozóan egy-egy kérdés szerepel a kérdésblokk első és második felében. Ezen túlmenően az eljárás

azt is garantálja, hogy az egyes doménekhez kapcsolódó tételek majdnem egyenletesen oszlanak el a sorozatban. Erre azért van szükség, mert bizonyos esetekben az azonos doménre vonatkozó tételek jelentésben igen közel állnak egymáshoz: ha pedig hasonló jelentésű tételek szorosan egymás után következnek, fennáll a veszélye, hogy a kitöltő nem dolgozza fel a mondatok jelentését kellő mélységben.

EREDMÉNYEK

Az öt személyiségdimenzióra vonatkozó összegzett mutatókat az egyes doménekhez tartozó tételekre kapott válaszok átlagértéke adta (természetesen az átlagok kiszámítását a fordított tételek tükrözését követően végeztük el). Validált mérőeszköz lévén, várakozásainknak megfelelően mind az öt mutató eloszlása halom alakú és közelítően szimmetrikus.

Az öt személyiségdimenzió és a négy védelemmotivációs tényező közötti nyers korrelációk meglehetősen összetett mintázatot mutatnak, részben azért, mert az öt személyiségdimenzió nem független egymástól, hanem egymással is bonyolult korrelációs viszonyban állnak. E nulladrendű korrelációkat a 7.2. táblázat foglalja össze:

7.2. táblázat: A személyiségdomének és a védelemmotiváció összetevői közötti korrelációk (Pearson-féle együtthatók)

	Válaszhatékonyság	Önhatékonyság	Személyes veszélyeztetettség	A fenyegetés súlyossága
Nyitottság	-0,072	0,053	0,216**	0,065
Lelkiismeretesség	0,259**	0,356**	-0,101	-0,019
Extraverzió	0,113	0,236**	0,059	0,059
Barátságosság	0,069	0,089	0,152*	0,209**
Neuroticizmus	-0,171*	-0,193**	0,180**	0,190**

Megjegyzés: * $p < 0,05$; ** $p < 0,01$

Forrás: saját szerkesztés

Az eredményekből jól látszik, hogy a lelkiismeretesség van a legerősebb hatással a védelemmotivációra, és e hatás épp a hipotetikus kauzális modellünk kulcsfontosságú első szakaszában, azon belül is az aktív védekező viselkedések megjelenéséért elsősorban felelős önhatékonyság esetében érvényesül. Vagyis elmondható, hogy az egyébként is lelkiismeretes, megbízható, szorgalmas, precíz, fegyelmezett emberekben nagyobb valószínűséggel alakul ki a meggyőződés, hogy lehetséges hatékonyan védekezni, és ők maguk képesek is megfelelően védekezni az online fenyegetések ellen (feltehetően annak köszönhetően, hogy valóban veszik is a fáradságot ahhoz, hogy megismerjék és elsajátítsák e megküzdési stratégiákat).

Az önhatékonysággal az extraverzió is pozitívan korrelál, bár e hatás nehezebben magyarázható a lelkiismeretesség hatásánál. Nem zárható ki ugyanis, hogy az önhatékonyságra vonatkozó kérdések esetében megmutatkozó pozitív énkép nem valódi képességeknek, hanem az extrovertált, beszédes, társaságkedvelő, aktív, lelkes személyekkel egyébként is összefüggésbe hozható önbizalom kifejeződése.

Érdekes mintázatot mutat továbbá a neuroticitás, amely a védelemmotiváció valamennyi komponensével szignifikánsan korrelál, és a korrelációk iránya (előjele) minden esetben a várt összefüggést mutatja. A neurotikus, érzelmileg labilisabb, aggodalmaskodó, szorongó egyének gyakrabban vélekednek úgy, hogy (1) az online fenyegetések ellen nem lehet hatékonyan védekezni, (2) ők maguk sem képesek erre, (3) valószínű, hogy a közeljövőben áldozatává válnak egy ilyen visszaélésnek, és (4) ez rendkívül súlyosan fogja érinteni őket.

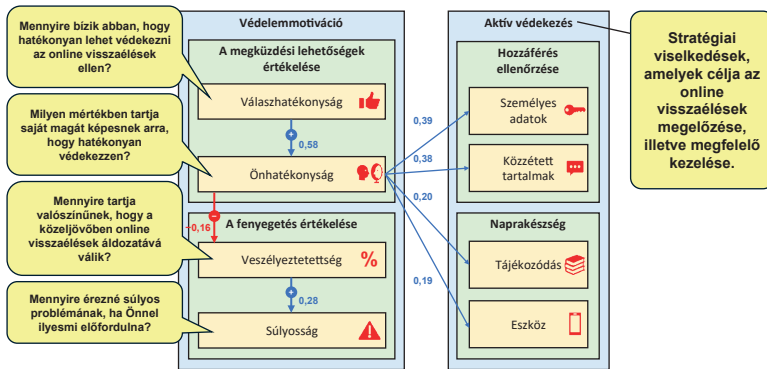
Ez utóbbi két vélekedés a barátságossággal is szignifikánsan korrelál, vagyis a barátságos, kedves, udvarias, jó természetű, szerény emberek negatívakban ítélik-értékelik az online világ fenyegetéseit.

Végül az új élményekre nyitott, ingerkereső, kíváncsi, érdeklődő személyek valószínűbbnek tartják, hogy a közeljövőben valamilyen visszaélés áldozatává válnak. Erre az összefüggésre magyarázatként kínálkozik, hogy a kíváncsi természetű emberek szélesebb körben tájékozódnak, így több hírről értesülnek ilyen eseményekkel kapcsolatban. Ez a magyarázat azonban

felveti a kérdést, hogy ugyanezekhez a személyekhez ugyanakkor a jelek szerint miért nem jutnak el nagyobb arányban az információk a védekezés különféle módjainak hatékonyságáról.

Ahogy korábban említettük, az egyszerű korrelációelemzés eredményeinek értelmezését megnehezítheti, hogy maguk a személyiségdimenziók sem tisztán ortogonális (azaz: egymástól tökéletesen független) faktorok, hanem egymással is gyengébb-erősebb korrelációs összefüggéseket mutatnak. E hatásoktól megszabadulhatunk, és némileg tisztább képet nyerhetünk a személyiség szerepéről a kiberbiztonsági védelemmotiváció elemeinek kialakulásában, ha egyszerű korrelációk helyett többváltozós regressziós modellekkel kísérreljük megjósolni a védelemmotiváció egyes tényezőit. Ha a regressziós modellekben az öt személyiségdimenzió egyszerre szerepel potenciális prediktorként, akkor az egyes dimenziókra vonatkozó statisztikai próbák azt mutatják meg, hozzájárul-e az adott dimenzió unikális módon a modell prediktív erejéhez. Vagyis ez esetben elemzéseink arra adnak választ, hogy ennek a személyiségdimenziónak van-e további magyarázó ereje azon felül, ami a másik négy dimenzió együttes figyelembevételével már meghatározható. E regresszióelemzések azt mutatják, hogy a válaszhatékonyosság a lelkiismeretességgel függ össze szignifikánsan és specifikusan ($\beta = 0,21, p = 0,009$), az önhatékonyosság szintén a lelkiismeretességgel ($\beta = 0,30, p < 0,001$), a személyes veszélyeztetettség érzete a nyitottsággal ($\beta = 0,17, p = 0,013$), a fenyegetés súlyosság érzete pedig a barátságossággal ($\beta = 0,20, p = 0,005$) és a neuroticizmussal ($\beta = 0,22, p = 0,005$) egyaránt.

A 7.1. ábra a 6. fejezetben felvázolt modell kiegészítése a személyiségdimenziók hatásaival. Ezen az ábrán kizárólag az összes többi személyiségdimenzió kontrollja mellett is megmutatkozó, unikális hatásokat tüntettük fel, és a nyilak mellett a standardizált béta-együtthatók szerepelnek.



7.1. ábra: A személyiség, a védelemmotiváció és az aktív védekezés dimenzióinak összefüggérendszer – a személyiség és a védelemmotiváció összetevői közötti kapcsolatok erősségét a többszörös regressziós modellek béta-együtthatói mutatják; az összes többi összefüggés esetében a nulladrendű korrelációs együtthatót tüntettük fel
 Forrás: saját szerkesztés

ÖSSZEGZÉS

Az eredmények alapján a Nagy Ötök személyiségmodell dimenziói szignifikáns hatást gyakorolnak az online fenyegetésekkel szembeni védelemmotivációra, megerősítve, hogy a személyiségjegyek a védelemmotiváció különböző komponenseinek felerősödésére vagy gyengülésére hajlamosítanak. Az alábbiakban az eredményeket a Nagy Ötök modell dimenzióival vetjük össze, és értelmezzük az összefüggéseket.

Lelkiismeretesség és önhatékonyság

Az eredmények szerint a lelkiismeretesség az önhatékonytságot, ezáltal pedig az aktív védekező viselkedések kialakulását erőteljesen befolyásolja. Ez összhangban áll a Nagy Ötök modell elméletével, amely szerint a lelkiismeretesség dimenziójába tartozó egyének szorgalmasak, felelősségteljesek és szervezettek, így hajlamosabbak erőfeszítéseket tenni a veszélyek megértésére és azok elhárítására. A lelkiismeretesség hatása azért különösen kiemelkedő, mert az önhatékonytságot erősítve meggyőződést formál arról, hogy a fenyegetésekkel szembeni védekezés nemcsak lehetséges, hanem hatékony is. Ez megerősíti McCrae és Costa állítását,⁷⁴ miszerint a lelkiismeretesség közvetlenül összefügg a célorientált és szabálykövető viselkedéssel.

Extraverzió és önhatékonyság

Az extraverzió pozitív korrelációja az önhatékonytsággal szintén jelentős, bár nehezebben értelmezhető. Az eredmények azt sugallják, hogy az extravertált személyek önbizalma hozzájárulhat ahhoz, hogy magasabb szintű önhatékonytságról számoljanak be, függetlenül attól, hogy ez tényleges képességeken alapul-e. Ez egybevág azzal az általános megfigyeléssel, hogy az extravertált egyének pozitív énképük és szociális készségeik miatt gyakran magasabbra értékelik képességeiket.⁷⁵ Ugyanakkor az extraverzió kockázati tényezőként is értelmezhető, mivel az online térben gyakran magasabb szociális aktivitás figyelhető meg az extravertált személyeknél, ami nagyobb kitettséget eredményezhet.

Neuroticizmus és védelemmotiváció

A neuroticizmus szignifikáns negatív hatása az önhatékonytságra és az általános védelemmotivációra összhangban áll az elméleti várakozásokkal.

⁷⁴ McCRAE–COSTA 1997.

⁷⁵ McELROY et al. 2007.

Az érzelmileg instabil, szorongó egyének hajlamosak eltúlozni a fenyegetések súlyosságát, miközben alábecsülik saját védekezési képességeiket. Ez a minta rávilágít arra, hogy a neurotikus személyek számára a fenyegetettségészlelés sokkal intenzívebb, ami ugyanakkor bénító hatással lehet a védekező viselkedés kialakulására. Az ilyen egyének gyakran érzik úgy, hogy tehetetlenek a veszélyekkel szemben, ami elkerülő vagy passzív magatartáshoz vezethet.⁷⁶

Barátságosság és fenyegetettségészlelés

Az eredmények szerint a barátságosság negatívan korrelál a fenyegetettségészleléssel, ami arra utal, hogy a barátságos, empatikus személyek kevésbé érzékelik az online fenyegetéseket súlyosnak. Ez összhangban van a Nagy Ötök modellel azon megállapításával, miszerint a barátságosság a másokba vetett bizalommal és a konfliktusok elkerülésével társul.⁷⁷ A barátságos egyének hajlamosak a jóindulatú értelmezésre, ami online kontextusban csökkentheti a védekezési hajlandóságot.

Nyitottság az élményekre és fenyegetettségészlelés

Az új élményekre nyitott személyek esetében az eredmények érdekes kettősséget mutatnak. Bár ezek az egyének valószínűbbnek tartják, hogy valamilyen visszaélés áldozatává válnak, ez az észlelés nem társul magasabb önhatékonysággal vagy védekező hajlandósággal. Ez arra utalhat, hogy bár a kíváncsi természetű egyének szélesebb körben értesülnek az online fenyegetésekről, az információk elsősorban a fenyegetettségük érzékelését növelik, nem pedig a védekezés hatékonyságába vetett hitüket. Ez a minta felveti annak szükségességét, hogy a kíváncsi egyének számára célzott információs kampányok révén ne csak a fenyegetések súlyosságát, hanem a védekezési lehetőségek elérhetőségét és hatékonyságát is hangsúlyozzák.

⁷⁶ ROGERS 1975.

⁷⁷ MCCRAE–COSTA 1997.

Javaslatok

Az eredmények jól mutatják, hogy a Nagy Ötök személyiségdimenziói közvetett módon, a védelemmotiváció különböző komponensein keresztül hatnak az online védekező viselkedések kialakulására. A lelkiismeretesség kiemelkedő szerepe az önhatékonyság kialakításában rávilágít arra, hogy a célzott oktatási programok és tudatosságnövelő kampányok különösen hatékonyak lehetnek ebben a csoportban.

Ugyanakkor az extravertált és neurotikus személyek esetében az önhatékonyság érzékelt szintje és a valós védekezési képességek közötti eltérés kezelése fontos kihívás. Ez megköveteli, hogy a kiberbiztonsági oktatás ne csak a technikai készségek fejlesztésére fókuszáljon, hanem az érzelmi reakciókat is vegye figyelembe, és reális önértékelésre ösztönözze az egyéneket.

A nyitottság és a barátságosság esetében a kommunikációs stratégiáknak ki kell egyensúlyozniuk a fenyegetettségészlelést és az önhatékonytágot. Ezek az eredmények megerősítik, hogy a személyre szabott intervenciók és képzési programok fejlesztése elengedhetetlen az online fenyegetésekkel szembeni védekezés hatékonyságának növeléséhez. A jövőbeli kutatásoknak érdemes lenne vizsgálni, hogy miként lehet a személyiségdimenziók hatását közvetlenül a védekező viselkedésekre optimalizálni.

8. A MESTERSÉGES INTELLIGENCIA ÁLTAL GENERÁLT KÉPEK FELISMERÉSÉNEK KÉPESSÉGE

A képek az emberi kommunikáció és a történetmesélés alapvető részét képezik.⁷⁸ A vizuális információk az üzenetek, érzelmek gyors és hatékony közvetítői, és jelentős mértékben befolyásolják a befogadó gondolkodását, attitűdjeit és viselkedését. A képek ereje abban rejlik, hogy képesek megragadni a figyelmet, érzelmi reakciókat kiváltani és emlékezetünkbe vésődni.⁷⁹ A vizuális elemek jelentősége a médiafogyasztásban is megmutatkozik, hiszen a hírek, cikkek és online tartalmak befogadását, értelmezését és megosztását is nagymértékben befolyásolják.⁸⁰ A vizuális tartalmak érzelmi töltete különösen fontos szerepet játszik a befogadók meggyőzésében és a véleményük formálásában.⁸¹ A vizuális kommunikációnak köszönhetően összetett gondolatok, információk és narratívák is könnyen befogadhatók válnak, és hatékonyabban jutnak el a célközönséghez.

A képek nem csupán információkat közvetítenek, hanem kulturális értékeket, társadalmi normákat és esztétikai preferenciákat is hordoznak.⁸² A vizuális kultúra korában a képek értelmezése és befogadása összetett folyamat, amelyet számos tényező befolyásol, többek között a befogadó kulturális háttere, személyes tapasztalatai és kognitív képességei.⁸³ Ugyanakkor, bár a vizuális percepció kétségtelenül hozzájárul a valóság megismeréséhez, nem feltétlenül fedi fel annak teljességét, különösen a mediatisztált képi világban. A képek manipulációja, amelyet a modern számítógépes technológiák lehetővé tesznek, tovább árnyalja a „látni és hinni” közötti hagyományos

⁷⁸ CASAS–WILLIAMS 2018.

⁷⁹ POWELL et al. 2015.

⁸⁰ WANG et al. 2023.

⁸¹ PALETZ et al. 2023.

⁸² GRUBER–HAUGBOLLE 2013.

⁸³ WALKER–WALKER 2004.

összefüggést.⁸⁴ A fényképek-videók módosításától a speciális effektéken át a virtuális valóságig terjedő technikai lehetőségek arra kényszerítenek minket, hogy kritikai szemmel vizsgáljuk a média által közvetített vizuális információkat, hiszen azok hitelessége egyre inkább megkérdőjelezhető.

A képekbe vetett bizalom változása érdekes kettősséget mutat a szakirodalomban: Westling kutatása azt az álláspontot képviseli, hogy az utóbbi időben az okostelefonokon található szerkesztőapplikációk, valamint a beépített kép- és videófilterek miatt a képek igazságtartalma folyamatosan veszített hitelességéből.⁸⁵ Vele szemben Parry úgy véli, hogy ennek tudatában és ennek ellenére továbbra is hajlamosak vagyunk elhinni azt, amit látunk.⁸⁶ Az igazi kihívást mégis az jelenti, hogy az emberek az archoz kapcsolják a fényképek és a digitális médiaproduktumok igazságértékét, így a manipulált tartalmak is könnyebben fejtenek ki hatást.⁸⁷ A képek hitelességének kérdése tehát egyre komplexebb és problematikusabb a digitális korban, és ez a probléma kritikus méreteket ölt a mesterséges intelligencia térnyerésével. Az MI ugyanis nem csupán új eszközöket adott a képek manipulálásához, hanem radikálisan átalakította a vizuális kommunikációt.

A generatív hálózatok (*General Adversarial Networks*, GAN)⁸⁸ és a diffúziós modellek, mint amilyen a DALL-E 3, a Stable Diffusion vagy a Midjourney⁸⁹ révén az MI által generált képek minősége az elmúlt években jelentősen javult, és ma már szinte megkülönböztethetetlenek a valódi fényképektől.⁹⁰ Ez a fejlesztés nemcsak művészeti alkotások, hanem az alternatív valóságot leképező, hamis arcokat tartalmazó mozgóképes tartalmak, úgynevezett *deepfake* videók létrehozására is alkalmas.⁹¹ A technológia

⁸⁴ BERGER 2012.

⁸⁵ KIETZMANN et al. 2019.

⁸⁶ PARRY 2009.

⁸⁷ KIETZMANN et al. 2019.

⁸⁸ GOODFELLOW et al. 2020.

⁸⁹ DHARIWAL–NICHOL 2024; ROMBACH et al. 2022.

⁹⁰ NIGHTINGALE–FARID 2022.

⁹¹ RICKER et al. 2024.

demokratizálódásával – a széles körben elérhető, felhasználóbarát alkalmazásoknak köszönhetően – bárki könnyedén készíthet megtévesztő képeket, ami jelentősen megnöveli a potenciális áldozatok körét. A valóság-hű hamisítványok veszélye abban rejlik, hogy nemcsak a vizuális élményt, hanem a hozzá kapcsolódó emlékeket és hiedelmeket is torzíthatják.⁹² Az MI által generált képek felismerése nem csupán technikai kérdés, hanem az emberi észlelés és kognitív folyamatok megértését is igényli.⁹³ A technológia gyors fejlődése miatt azonban egyre nehezebb lépést tartani a legújabb módszerekkel és a hamisítványok egyre kifinomultabb formáival.⁹⁴

A manipulált képek felismerése komoly kihívást jelent a laikusok és a szakértők számára egyaránt. Több kutatás is igazolja (többek között: Veszelszki és szerzőtársai⁹⁵ vagy Nightingale és szerzőtársai⁹⁶), hogy az emberek nehezen tudják megkülönböztetni a valós és az MI által generált képeket, különösen akkor, ha a hamisítványok kiváló minőségűek és realisztikusak. A manipuláció felismerését számos tényező nehezítheti: többek között a képek komplexitása, a befogadó előzetes tudása és tapasztalata, valamint az érzelmi reakciók, amelyeket a képek kiváltanak. Az MI által generált képek esetében gyakran előfordulnak apróbb hibák és anomáliák, amelyek árulkodó jelek lehetnek a szakértők számára.⁹⁷ Ezek lehetnek például szokatlan textúrák, aránytalan arcvonások, természetellenes fények és árnyékok vagy a kép kontextusával nem összeegyeztethető elemek.⁹⁸ A manipulált képek felismerésében az emberi intuíció is szerepet játszhat, azonban ez nem mindig megbízható, és további vizsgálatokat igényel.⁹⁹ A *deepfake* technológia fejlődése és a generatív MI-modellek egyre kifinomultabb képességei miatt

⁹² NEWMAN et al. 2012.

⁹³ LU et al. 2023.

⁹⁴ MARTIN-RODRIGUEZ – GARCIA-MOJON – FERNANDEZ-BARCIELA 2023.

⁹⁵ VESZELSZKI–HORVÁTH–KOVÁCS 2022.

⁹⁶ NIGHTINGALE–WADE–WATSON 2017.

⁹⁷ YU et al. 2024.

⁹⁸ WANG et al. 2020.

⁹⁹ FRANK et al. 2024.

a manipulált képek felismerése folyamatos kihívást jelent, és további kutatásokra van szükség a hatékonyabb detektálási módszerek kidolgozásához.¹⁰⁰ A kutatások azt mutatják, hogy az emberek hajlamosak bízni a látott képekben és videóban, még akkor is, ha azok manipuláltak.¹⁰¹ Ezért elengedhetetlen, hogy a médiahasználók tisztában legyenek a *deepfake* technológia veszélyeivel, és képesek legyenek felismerni a hamisítványokat.

Kutatásunk célja annak vizsgálata, hogy a jelenlegi fejlettségi szinten mennyire tudják a laikus médiahasználók megkülönböztetni az MI által generált képeket a valódi fényképektől. Fontos megvizsgálni, hogy mekkora eltérések mutatkoznak az emberek teljesítményében, és milyen tényezők befolyásolják a sikeres detektálást. Ezen kérdések megválaszolása kulcsfontosságú a médiaismeret fejlesztéséhez, a manipuláció elleni küzdelemhez és a digitális térben való biztonságos eligazodáshoz.¹⁰²

MÓDSZER

Kutatásunk során arra voltunk kíváncsiak, milyen tényezők befolyásolják a kitöltőket a mesterséges intelligencia által generált és valódi képek felismerésének képességét illetően. A folyamat első lépéseként meghatároztuk a valódi képeket, és azokat alapul véve promptoltuk a Playground AI nevet viselő képgenerálásra használható platformot. A képeket 2024 tavaszán generáltuk, így a platform akkori képességeit tükrözik. Ezt fontos figyelembe venni, hiszen a generatív mesterséges intelligencia modellek folyamatosan fejlődnek, és a képességeik egyre meggyőzőbbek.

A felhasznált valódi képeket ingyenes stockfotóbázisból választottuk ki. Mind férfiakat, mind nőket ábrázoló képekből 3-3 portrét választottunk azonos szempontok szerint: a képeken hozzátétőlegesen 20-as, 40-es és 60-as

¹⁰⁰ LAĐEVIĆ et al. 2024.

¹⁰¹ NIGHTINGALE–WADE–WATSON 2017.

¹⁰² VESZELSZKI–HORVÁTH–KOVÁCS 2022.

éveikben járó férfiak és nők szerepeljenek. A mesterséges intelligenciával generált képek esetében figyeltünk arra, hogy a mesterséges intelligencia promptolása során azonos technikai specifikációkat adjunk meg, mint amivel az eredeti képeket készítette az alkotója. A generálás során a Playground AI fizetős verzióját használtuk, ugyanis ebben a formában sokkal több beállítást tudtunk alkalmazni, az ingyenes esetében a filter, különböző motorok változtatása nem állt volna módunkban. A függelékben található promptok segítségével készült képek közül végül közösen választottuk ki a felhasznált képeket.

A nyolcas csoportokból minden esetben egyetlen képet választottunk ki, azt, amelyet három ítéző közmegegyezéssel a legvalószerűbbnek gondolt. E kérdőívblokkban a képek sorrendjét véletlenszerűen, meghatározott korlátozások mentén adtuk meg. A végleges sorrend kialakításakor további fontos szempont volt, hogy csökkentjük annak az előfordulási valószínűségét, hogy a kitöltő hozzászokik egy bizonyos ingertípushoz: ha például egymás után túl sok mesterséges intelligencia által generált képet lát, ez sorozathatásként befolyásolhatja a képek percepciójának folyamatát. A végső sorrendet egy olyan algoritmussal állítottuk elő, amely a képek véletlenszerű permutációit állította elő, ezt ismételve addig, amíg nem talált egy olyan sorrendet, amely az alábbi feltételek mindegyikének megfelelt:

1. Azonos nemű személyből legfeljebb három követheti egymást a sorozatban.
2. Azonos korcsoportba tartozó személyekből legfeljebb kettő követheti egymást.
3. Azonos típusba tartozó ingerből (valódi fénykép, illetve MI által generált kép) legfeljebb három követheti egymást.

A végleges bemutatási sorrend a fentiek értelmében a 8.1. táblázatban látható elrendezés szerint alakult.

8.1. táblázat: A képi ingerek jellemzői és bemutatási sorrendje

Pozíció	ID	Nem	Életkor	Típus
1.	1	férfi	fiatal	valódi
2.	12	nő	idős	MI-generált
3.	5	nő	középkorú	valódi
4.	2	férfi	középkorú	valódi
5.	10	nő	fiatal	MI-generált
6.	9	férfi	idős	MI-generált
7.	11	nő	középkorú	MI-generált
8.	4	nő	fiatal	valódi
9.	3	férfi	idős	valódi
10.	7	férfi	fiatal	MI-generált
11.	8	férfi	középkorú	MI-generált
12.	6	nő	idős	valódi

Forrás: saját szerkesztés

A vizsgálatban használt teljes képi ingeranyagot miniatűrök formájában a 8.2. táblázat foglalja össze. A képeket nagy méretben (ahol a részletek is megfigyelhetők) a kötet függeléke tartalmazza.

8.2. táblázat: A vizsgálatban használt valódi fényképek és mesterséges intelligencia által generált képi ingerek

		fiatal	középkorú	idős
Valódi fénykép	férfi			
	nő			
MI által generált kép	férfi			
	nő			

Forrás: saját szerkesztés

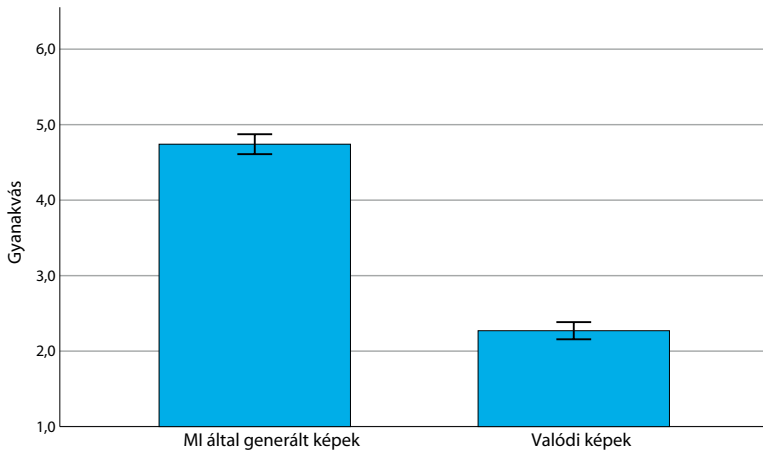
A kitöltőket nem szorítottuk időkeretek közé, így addig vizsgálhatták a képeket, amíg szükségesnek ítélték. Minden esetben 6 válaszból álló Likert-skálás kérdést tettünk fel a látott kép megítélésével kapcsolatban:

- ♦ Szinte biztos, hogy ez valódi fénykép.
- ♦ Valószínű, hogy ez valódi fénykép.
- ♦ Bizonytalan vagyok, de inkább azt mondanám, hogy ez valódi fénykép.
- ♦ Bizonytalan vagyok, de inkább azt mondanám, hogy ezt a képet MI generálta.
- ♦ Valószínű, hogy ezt a képet MI generálta.
- ♦ Szinte biztos, hogy ezt a képet MI generálta.

A kitöltő válasza tehát mind a valódi fényképek, mind az MI által generált képek esetében egységesen az MI-re való gyanakvás fokát jelöli. Amennyiben az adatközlő nem a „Szinte biztos, hogy ez valódi fénykép” választ jelölte, arra kértük, hogy röviden írja le, mi miatt érzi úgy, hogy a fénykép nem biztos, hogy valódi. A képnek mely elemei tűntek gyanúsak vagy egyenesen áruklodónak?

EREDMÉNYEK

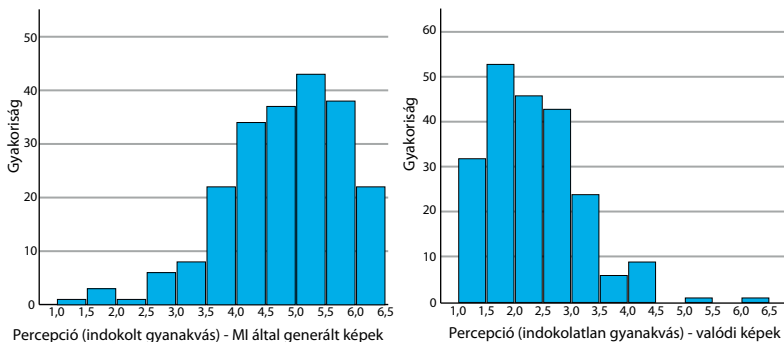
Az eredményeink összességében azt mutatják, hogy a kitöltők nagy biztonsággal képesek megkülönböztetni a valódi fényképeket a jelenleg széles körben hozzáférhető MI-eszközök által generált képektől (a felhasználók számára széles körben hozzáférhető technológia mai fejlettségi szintjét vizsgálatunkban a Playground AI platform előfizetős verziója reprezentálta). A hatos skálán a gyanakvás átlagos szintje MI-képek esetében 4,74 (szórás: 0,98), míg valódi képek esetében az (indokolatlan) gyanakvás átlagértéke 2,27 volt (szórás: 0,85). Az illesztett mintás t -próba eredménye szerint a különbség erősen szignifikáns: $t(214) = 28,433$; $p < 0,001$. A két ingertípus esetében a gyanakvás szintje közötti korreláció nem szignifikáns, mértéke elhanyagolható ($r = 0,034$), tehát az adatsorban nincs olyan mintázat, amely arra utalna, hogy a képek megítélését befolyásolná a kitöltő valamiféle általános hajlama a gyanakvásra, a kép jellegétől függetlenül. A két ingertípus megítélésében mutatkozó igen nagy mértékű különbséget (Cohen-féle $d = 1,939$) 95%-os konfidenciaintervallumokkal a 8.1. ábra szemlélteti.



8.1. ábra: Az MI által generált képek és valódi fényképek percepciója között mutatkozó különbség

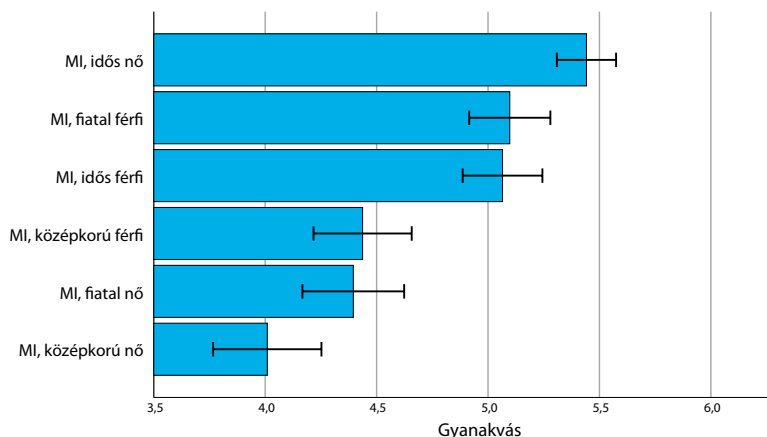
Forrás: saját szerkesztés

A két ingertípusra adott válaszok különbözősége a gyanakvás eloszlásán is jól megfigyelhető (8.2. ábra). Az MI által generált ingerek esetében balra ferde, míg a valódi képek esetében jobbra ferde eloszlást kaptunk, amelyet a 8.2. ábra hisztogramjai is megerősítenek.



8.2. ábra: A gyanakvás mértékének eloszlása MI által generált képek és valódi fényképek esetében

Forrás: saját szerkesztés

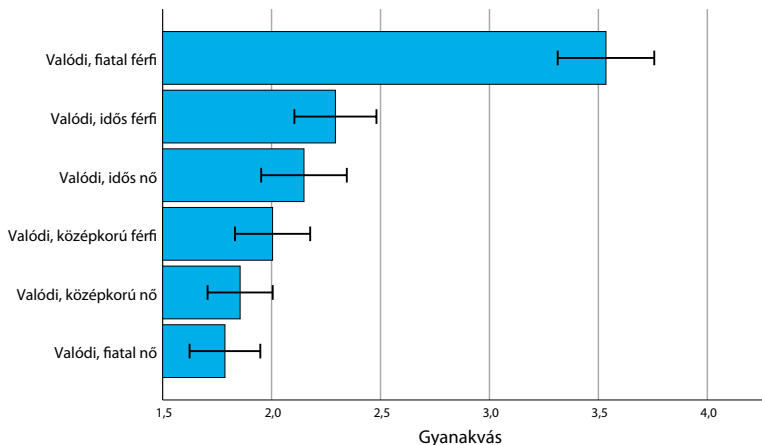


8.3. ábra: A gyanakvás átlagos mértéke a hat MI által generált kép esetében

Forrás: saját szerkesztés

Az MI által generált képek sikeres detektálása (vagyis az indokolt gyanakvás mértéke) természetesen ingerenként különbözött. A kitöltők leginkább az idős nőt ábrázoló képet vélték gyanúsnak (átlag: 5,44), de hasonlóan könnyen ismerték fel az MI jellegzetességeit a fiatal férfit és az idős férfit ábrázoló képeken is (átlagok, rendre: 5,10 és 5,07). A legnagyobb kihívást a középkorú nőt ábrázoló kép jelentette, itt az átlagos gyanakvás mértéke alig haladta meg a skála 3,5-ös középvértékét (átlag: 4,01). Általánosan elmondható ugyanakkor, hogy valamennyi MI-generált inger esetében a válaszadók többsége arra hajlott, hogy a kép valószínűleg nem valódi. A hat ingerhez kapcsolódó gyanakvás átlagos szintjét 95%-os konfidenciaintervallumokkal a 8.3. ábra mutatja.

A gyanakvás azonban nem kizárólag az MI által generált ingerek esetében érdekes jelenség, hiszen e technológia megjelenésével megjelent a másik irányú percepció torzítás is, amely azt eredményezi, hogy bizonyos felhasználók már a valódi fényképeken is látni vélik a MI-algoritmusok tipikus jellemzőit, azaz a tényleges fényképek valódiságát is megkérdőjelezzik. Esetünkben ennek az indokolatlan gyanakvásnak a szintje – mint fent láttuk – alacsony (átlag: 2,27). Az egyes képekre külön lebontva azonban



8.4. ábra: A gyanakvás átlagos mértéke a hat valódi fénykép esetében

Forrás: saját szerkesztés

azt látjuk, hogy ez nem minden esetben van egységesen így. Esetünkben a valódi fiatal férfiről készült képpel kapcsolatban érezték a kitöltők leginkább, hogy azt MI generálhatta (a gyanakvás átlagértéke: 3,53), míg a többi kép esetében igen magabiztosan állították, hogy a kép valódi (átlagok: 1,79–2,29). E mintázatot – 95%-os konfidenciaintervallumokkal – a 8.4. ábra szemlélteti.

A valódi és az MI által generált képekre külön-külön számított gyanakvási mutató mellett minden kitöltőre kiszámítottunk egy összesített mutatót is, amely a percepció pontosságát mutatta. E mutató értékét a következő képlettel határoztuk meg:

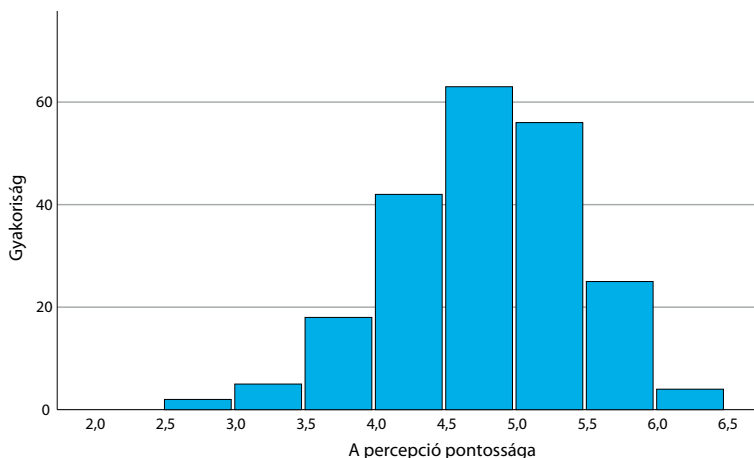
$$P = \frac{GY_{MI} + (7 - GY_v)}{2}$$

ahol P a percepció pontosság, GY_{MI} a MI-képek esetében mutatott, indokolt gyanakvás átlagos szintjét, a GY_v pedig a valódi képek esetében mutatott, indokolatlan gyanakvás átlagos szintjét jelöli. A mutató tehát az elméletileg legmagasabb 6-os értéket akkor éri el, ha a kitöltő valamennyi MI-val generált kép esetében a „Szinte biztos, hogy ezt a képet MI generálta” opciót, és valamennyi valódi kép esetében a „Szinte biztos, hogy ez valódi fénykép”

opciót választotta. Az elméletileg legalacsonyabb, 1-es értéket az a kitöltő kapná, aki éppen fordítva, a MI-képek valóságáról és a valódi képek MI-generált voltáról lenne minden esetben kivétel nélkül meggyőződve.

E percepció mutató minimuma 2,75, míg maximuma 6,00 volt; tehát volt a mintánkban olyan adatközlő, aki mind a 12 ingerre nemcsak helyes választ adott, hanem teljes magabiztossággal tette azt. A percepció pontosságának átlagértéke 4,74, míg a szórása 0,64 volt. A további elemzés azt mutatta, hogy a mutató eloszlása halom alakú, közel szimmetrikus, bár kissé balra ferde, amely a skála középértékét meghaladó átlagértéknek köszönhető. Ezt szemlélteti a 8.5. ábrán látható hisztogram.

Ez az összesített mutató lehetővé tette azt is, hogy megvizsgáljuk, kimutathatók-e szignifikáns különbségek a percepció pontosságában különféle ingertípusok esetében. Az adataink szignifikáns eltéréseket mutatnak mind az ábrázolt személy neme, mind az ábrázolt személy korcsoportja tekintetében.

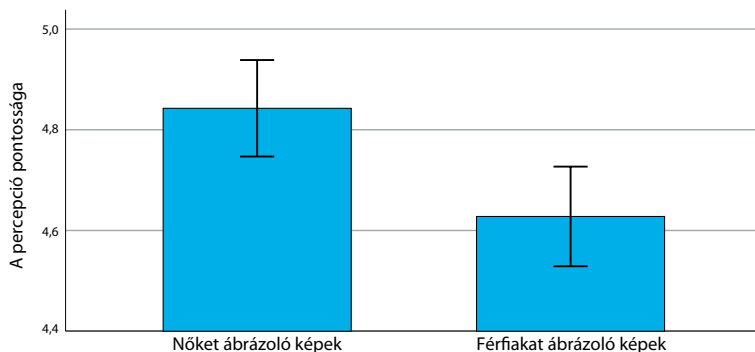


8.5. ábra: A percepció pontosságát mérő összesített mutató eloszlása

Forrás: saját szerkesztés

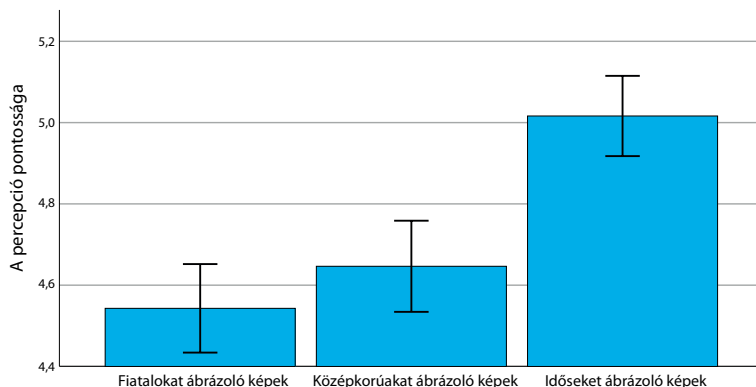
Az illesztett t -próba eredményei szerint a felhasználók szignifikánsan pontosabban ismerték fel a nőket ábrázoló, mint a férfiakat ábrázoló képeket: a pontosság átlagértéke a két esetben: 4,83, illetve 4,62. Az eltérés mértéke tehát kismértékű (Cohen-féle $d = 0,311$), viszont erősen szignifikáns: $t(214) = 4,567, p < 0,001$. A két ingertípus felismerésében mutatkozó különbséget – 95%-os konfidenciaintervallumokkal – a 8.6. ábra szemlélteti.

Hasonlóan szignifikáns hatású a percepció pontosságára az ábrázolt személy életkora. Az eredmények szerint az ismételt mérés ANOVA fontos előfeltétele, a szfericitás nem sérült (Mauchly-féle $W = 0,995; \chi^2[2] = 1,045$, n. s.), vagyis a feltételpáronkénti különbségek varianciája a populációban egyenlőnek tekinthető. Az ismételt mérésekkel végzett varianciaanalízis eredményei szerint – a szfericitás feltételezését megtartva – szignifikáns különbségek voltak a különböző korcsoportokat ábrázoló képek helyes felismerésében: $F(2, 428) = 39,498, p < 0,001$. Ahogy a 8.7. ábrán is jól látható, e különbségeket az okozza, hogy idős emberekről készített valódi és MI-generált képeket a kitöltők pontosabban tudták megkülönböztetni (átlagos pontosság: 5,02), mint a középkorú személyekről (átlag: 4,65), illetve a fiatal személyekről (átlag: 4,54) készített kétféle ingertípust. A 8.7. ábrán a hibahatárok – a korábbi ábrákhoz hasonlóan – 95%-os konfidenciaintervallumokat jelölnek.



8.6. ábra: A percepció pontossága a képen látható személy neme szerint

Forrás: saját szerkesztés



8.7. ábra: A percepció pontossága a képen látható személy életkora szerint

Forrás: saját szerkesztés

A gyanakvás szöveges indoklásában azonosítható mintázatok

Érdekes a statisztikai eredményeket összevetni a szöveges válaszokkal. A 8.8. ábrán látható összehasonlító szófelhőben jól látható, hogy az MI által generált képek esetében a gyanakvást legjellemzőbben valamiféle túlzás váltotta ki („túl”, „túlságosan”). E határozószókhoz gyakran kapcsolódtak a túl „tökéletes”, „szabályos”, „sima” melléknevek, amelyek jellemzően a „bőr”, „arc”, „haj”, „szem”, „kézfej” főnevek jelzőiként szerepeltek. Az ábra alsó felén, a valódi fényképekre adott válaszokban gyakran előforduló szavak körében nem figyelhető meg ilyen világos mintázat. A „férfi” szó például csak azért szerepel gyakran, mert a kép, amely a legtöbb adatközlőben váltott ki indokolatlan gyanakvást, egy fiatal férfit ábrázolt. A többi gyakori szó is azt jelzi, hogy a valódi fényképek esetében a gyanakvást egy általános benyomás („megérzés”) váltja ki, amely szerint a kép nem tűnik „természetes”-nek vagy „valódi”-nak, és ezt a (hamis) érzetet nagyjából hasonlóan gyakran váltja ki a „környezet” („fények”, „árnyékok”, „szoba”), illetve a képen látható élettelen tárgyak („függöny”, „ing”), mint a képen szereplő személy részletei („fogak”, „száj”, „kezek”).



8.8. ábra: Az MI-generált képek és a valódi fényképek által kiváltott gyanakvás indoklásában szereplő szavak gyakoriságát összehasonlító szófelhő

Megjegyzés: Az ábra felső részén olvasható zöld szavak az MI-képre adott válaszokban, míg az alsó részen látható narancssárga szavak a valódi fényképekre adott válaszokban fordultak elő gyakrabban. A szavak mérete a kétféle ingertípusra adott válaszokban a gyakoriságbeli különbséget fejezi ki.

Forrás: saját szerkesztés

A textuális adatok közelebbi vizsgálata alapján kirajzolódik, hogy bár az MI által generált képek egyre realisztikusabbak, a válaszadók mégis számos olyan vizuális jelle mutattak rá, amelyek segíthetnek a valódi és a mesterségesen

generált képek megkülönböztetésében – és nem mellesleg teljesen megfelelnek a szakirodalom által említett tényezőknek. Ebben, úgy látszik, nem sikerült olyan mértékű előrelépést megtennie az algoritmusok fejlesztőinek, ami kiküszöbölne a problémákat.

Az egyik leggyakoribb dilemmát a realizmus paradoxona okozta. Érdekes módon a „túl tökéletes” megjelenés, ami a korábbi MI-generációk gyengesége volt, továbbra is árulkodó jel volt a kitöltők számára, gyakran emlegették a bőr természetellenes simaságát („semmi pigmentáció, ránc, bőrhiba”), a hiányzó ráncokat, a túl élénk színeket és a hibátlan, „műanyag” hatású felületeket („minden túl sima, hiányoznak a textúrák a képről”). Ez arra utal, hogy miközben az algoritmus a realizmusra törekszik, olykor túllő a célon, és olyan képeket hoz létre, amelyekről hiányoznak a valós világ apró tökéletlenségei.

Voltak anatómiai buktatók, különösen a kezek és az ujjak ábrázolása jelent továbbra is némi kihívást az MI számára. A válaszadók számos példát hoztak a furcsa ujjtartásra, a hiányzó („4 ujj van csak”) vagy deformált („ráncatlanok az ujjai”) ujjakra, a természetellenes kéztartásra és a kézfej aránytalanságaira („Jobb kézfej térbeli elhelyezkedése természetellenes”; „Bal kézfej furcsán áll, hajlik”; „A kézfején lévő ráncok, bőrrödök, bőrszín”). A szemek és az arc is gyakran elárulják a mesterséges eredetet: a kifejezésten tekintet, a túl szabályos arcvonások, a szimmetria és a bőr textúrájának hiánya mind gyanús jelek voltak és lehetnek.

Fontos szerepet játszott a kontextus és a kompozíció, a háttér és a környezet elemzése, tagadhatatlan ennek jelentősége. A homályos, elnagyolt vagy természetellenes háttér („A háttér homályosítása alul körkörös” [sic!]), a fények és árnyékok furcsa eloszlása („nincsenek árnyékok, holott a fényviszonyok alapján kellene lenni”), a tárgyak irreális elhelyezkedése, valamint a személy és a környezet közötti diszharmónia („A háttér homályosítása alul körkörös, a test nagy része mögött pedig túl erős a *blur* technika”) mind utalhattak mesterséges eredetre. A ruházat és a kiegészítők is árulkodók voltak: a ruhák természetellenes esése, a kiegészítők elmosódottsága vagy torz formája, az értelmetlen feliratok és a kiolvashatatlan szövegek az MI-generált képeket jelentették.

A válaszadók bizonytalansága a valódi képek esetében is rámutat arra, hogy a modern képszerkesztő programok és filterek megnehezítik a valós és a generált képek megkülönböztetését. Ahogy az MI-technológia fejlődik, és egyre realisztikusabb képeket hoz létre, a detektálás is egyre nagyobb kihívást jelent. A kritikus gondolkodás, a részletekre való odafigyelés, az anatómiai ismeretek és a képi kompozíció elemzése mind fontos szerepet játszanak a MI által generált képek felismerésében. A jövőben feltehetően újabb módszerekre és eszközökre lesz szükség a mesterségesen generált képek azonosításához. A kutatásnak a technológiai fejlődéssel párhuzamosan kell haladnia, hogy megőrizhető (vagy inkább visszaállítható) legyen a vizuális információk hitelességébe vetett bizalom.

ÖSSZEGZÉS

Tanulmányunkban a mesterséges intelligencia által generált képek felismerésének képességét vizsgáltuk laikus médiahasználók körében. A Playground AI platform Stable Diffusion modelljével generált, realisztikus portrékat valódi fényképekkel vetettük össze egy online kérdőíves kísérletben. Eredményeink azt mutatják, hogy a válaszadók meglepő pontossággal tudták megkülönböztetni a kétféle képtípust. Az MI-generált képek esetében a gyanakvás átlagos szintje szignifikánsan magasabb volt (4,74), mint a valódi fényképek esetében (2,27). Ez a különbség igen nagy mértékűnek tekinthető (Cohen-féle $d = 1,939$). Fontos megjegyezni, hogy a technológia folyamatosan fejlődik, így a jövőben ez a különbség várhatóan csökkenni fog.

A percepció pontosságát vizsgáló összesített mutatónk alapján megállapítottuk, hogy a válaszadók átlagosan 4,74 pontot értek el a 6-os skálán, ami azt jelzi, hogy a legtöbb esetben helyesen azonosították a képek eredetét. Érdekes módon a nők képeit valamivel pontosabban tudták megkülönböztetni, mint a férfiakét, valamint az idősebb személyeket ábrázoló képek esetében volt a legmagasabb a percepció pontossága. Ez utóbbi eredmény összhangban van a szöveges válaszok elemzésével, amely szerint az MI-algoritmusok továbbra is nehézségekbe ütköznek az emberi anatómia, különösen a kezek

és az arc realiztikus ábrázolásában. A válaszadók gyakran említették a bőr természetellenes simaságát, a hiányzó ráncokat, a furcsa ujjtartást és a kifejeztelen tekintetet mint árulkodó jeleket. Emellett a háttér és a környezet elemzése is fontos szerepet játszott a képek megítélésében. A homályos, elnagyolt vagy természetellenes háttér, a fények és árnyékok furcsa eloszlása, valamint a személy és a környezet közötti diszharmonia mind utalhattak mesterséges eredetre.

Kutatásunk eredményei rávilágítanak a médiaismeret fejlesztésének fontosságára a digitális korban. A mesterséges intelligencia által generált képek egyre realiztikusabbak, és egyre nehezebb őket megkülönböztetni a valódi fényképektől. A kritikus gondolkodás, a részletekre való odafigyelés, az anatómiai ismeretek és a képi kompozíció elemzése mind fontos szerepet játszanak a manipulált képek felismerésében. Ezek most még olyan szinten állnak, hogy egy gyakorlott szem viszonylag könnyen kiszúrja őket, viszont érdekes kérdést vet fel ez a jövőre nézve: elvárható-e egy átlagos médiafogyasztótól, hogy szakértő legyen vizuális téren? A jövőben további kutatásokra van szükség a hatékonyabb detektálási módszerek kidolgozásához, valamint a médiahasználók felkészítéséhez a *deepfake* technológia veszélyeire is. Ahogy az algoritmusok fejlődnek, ezzel párhuzamosan a kutatásoknak is haladniuk kell, hogy megőrizhessük a vizuális információk hitelességébe vetett bizalmat, és megakadályozhassuk a manipulált képek általi visszaéléseket. Teljes pontossággal működő detektáló rendszerek hiányában ennyi az, amit tehetünk, bízva abban, hogy az egyéni teljesítmény összeadódik.

9. A VIZUÁLIS MI-FELISMERŐ KÉPESSÉG HÁTTERTÉNYEZŐI

A 2015-ben megrendezett World Press fotóversenyen huszonkét fotót zártak ki utólagosan beazonosított vizuális manipuláció miatt. A rangos verseny közfelháborodást váltott ki a fotóriporteri szakmában, és a botrányt követően a World Press új etikai szabályt hozott létre: egy képverifikáló teszt elvégzése mellett a fotósoknak bizonyítaniuk kell, hogy képeik tükrözik a valóságot, és nem vezetik félre a közönséget.¹⁰³ Ha egy világszinten ismert verseny szakmai zsűrije sem minden esetben képes kiszűrni az MI által generált vizuális tartalmakat, a laikus tartalomfogyasztóknak várhatóan sokkal kisebb esélyük van a helyes megítélésre.

A jelen kutatás bizonyítéku szolgál arra, hogy a bemutatott vizuális tartalmak jellemzői (például az ábrázolt személy neme és korcsoportja) erőteljesen befolyásolják a befogadó megítélését (lásd 8. fejezet). A valódi és MI által generált képek megkülönböztetésének képessége természetesen nemcsak az ingerek különféle jellemzőin múlhatnak, hanem a képet megtekintő felhasználó attribútumain is. Itt gondolhatunk a generációs különbségekre, az elsajátított digitális kompetenciákra vagy a felhasználók neméből adódó kognitív eltérésekre is. Az ebben a fejezetben bemutatott elemzések azt a kérdéskört vizsgálják, milyen egyéni jellemzők hozhatók összefüggésbe azzal, hogy valaki hatékonyabban vagy kevésbé hatékonyan képes MI által létrehozott vizuális tartalmakat megkülönböztetni a valódi fényképektől.

¹⁰³ NIGHTINGALE–WADE–WATSON 2017.

A DIGITÁLIS KÉPMANIPULÁCIÓ KORSZAKA

„Minél többet tudunk, annál többet látunk” – tartja az Aldous Huxley nevéhez fűződő közismert mondás.¹⁰⁴ Az infokommunikációs technológia kibővített meghatározása bőven túlmutat az okoseszköz-használton, különösen a közösségimédia-platformok által közvetített végtelen mennyiségű vizuális tartalom miatt.¹⁰⁵ A Photoshop, Affinity Photo és számos más képszerkesztő szoftver megjelenése új kapukat nyitott meg a digitális képszerkesztés számára, amit a mesterségesintelligencia-alapú, ingyenesen igénybevehető weboldalak, képszerkesztő applikációk követtek. Ezek a felületek nem csupán a professzionális fotográfusok számára elérhetőek – akik a laikusokkal ellentétben ismerik és követik az etikus képszerkesztés szabályait –, hanem a közösségi felületek bármely felhasználója használhatja ezeket, ezáltal filtert helyezhet a szelfijére, karcsúsíthatja a képen a derekát, vagy éppen létrehozhat egy MI-avatárt. A gyorsan letölthető, felhasználóbarát módon, könnyen kezelhető képszerkesztő-felületek felülírják „a kamera soha nem hazudik” kijelentést, ezáltal jogosan kérdőjelezhető meg a vizualitás és igazságtartalom összefüggése, vagyis egy képi produktum hitelessége.¹⁰⁶

Az új technológiák jelentősen megváltoztatták a képek, hangok és szavak társadalmunkra gyakorolt hatását. A digitális környezetünkben történt korszakos átszerveződés a valósághoz és az egymáshoz való viszonyulásunkat is meghatározza.¹⁰⁷ A megváltozott társadalmi elvárások új műveltség megjelenését segítették elő, amelyet vizuális műveltségnek hívunk.¹⁰⁸ A vizuális műveltség (Spalter és Van Dam alapján „digitális vizuális műveltség”-ként is ismert; *digital visual literacy*);¹⁰⁹ a különböző vizuális produktumok értelmezésének és elkészítésének képessége.¹¹⁰ Spalter és van Dam alapján három alapvető kompetencia határozza meg a digitális vizuális műveltséget:

¹⁰⁴ HUXLEY 1974.

¹⁰⁵ YOUSSEF–DAHMANI–RAGNI 2022.

¹⁰⁶ NEWTON 2006.

¹⁰⁷ GALLESE 2024.

¹⁰⁸ MESSARIS 2012.

¹⁰⁹ SPALTER – VAN DAM 2008.

¹¹⁰ DEBES 1969.

1. a digitális-vizuális produktumok kritikus értékelése (kétdimenziós, háromdimenziós, statikus és mozgó);
2. a digitális-vizuális adatok (etikus) ábrázolásáról való döntés;
3. technológiaalapú hatékony vizuális kommunikáció a befogadókkal.

AZ MI-FELISMERŐ KÉPESSÉG

Joggal feltételezhető, hogy a digitális médiával felnövő generációk vizuális műveltsége egyenesen arányos a hatékonyabb képi felismeréssel, viszont a képszerkesztési gyakorlatok ismerete és a vizuális manipuláció felismerésének képessége még mindig meglehetősen alacsony a digitális bennszülöttek körében. Ezt igazolja Lazard és kutatócsoportjának vizsgálódása is, amelyben szépségápolási hirdetések eltérő arányban manipulált képeinek fogyasztói attitűdre gyakorolt hatását vizsgálták egyetemi hallgatók körében.¹¹¹ Fő kutatási kérdésük az volt, vajon befolyásolja-e a képek manipulálási aránya a fogyasztói döntést. A hirdetésen minden esetben egy modell volt látható, akinek elsősorban az arcán végeztek szerkesztést: nagyobb szemek, fehérebb fogak, szimmetrikusabb arc és szebb bőr. A kutatás arra a következtetésre jutott, hogy a képek megtekintési idejétől függetlenül a résztvevők csak kismértékben tudták azonosítani a manipulált elemeket. Abban az esetben, amikor választaniuk kellett a nagyobb vagy kisebb mértékben manipulált képek közül, pozitívabb attitűddel rendelkeztek a nagyobb mértékben manipulált hirdetések esetében.

Egy másik kutatás a politikai *deepfake* videók vizuális hamisításának felismerését tanulmányozta 829 egyetemi hallgatóval.¹¹² A kutatás a képhamisítás kérdéskörét járta körül, és azt vizsgálta, milyen motivációk alapján képesek a befogadók megkülönböztetni egy valódi politikai beszédet a *deepfake* által generálttól. Öt fő okot emeltek ki, amelyek miatt a résztvevők bizalmatlanok voltak az üzenetekkel szemben: 1. az elhangzottak és a politikai realitás közötti eltérés (30,6%); 2. véleményen alapuló politikai bizalmatlanság

¹¹¹ LAZARD–BOCK–MACKERT 2020.

¹¹² HAMELEERS – VAN DER MEER – DOBBER 2024. Köszönjük a kutatás ajánlását Taxner Tünde kollégánknak.

(15,9%); 3. az audiovizuális ingerek észlelt manipulációja (12,1%); 4. tényyszerű állítások, források hiánya (10,3%); 5. médiatudatosság (2,7%). A résztvevők 28,5%-a egyáltalán nem tudott választ adni az audiovizuális manipuláció beazonosítására, emellett pedig az alanyok mindössze 12,1%-a vette észre a videók közötti tényleges különbségeket. A résztvevők főként a szintetikus szájmozgást emelték ki hamis elemként. Ám a felmérés fő konklúziója az volt, hogy a befogadók általánosságban hitelesnek vélik a *deepfake* videókat.

Ugyancsak a politikai kommunikáció területéről származik Dobber és kutatótársai¹¹³ vizsgálata. A politikai attitűdök *deepfake*-kel történő megváltoztatására irányuló online kísérletükben azt találták, hogy az általuk létrehozott *deepfake*-videó képes volt egy politikai szereplő megítélését megváltoztatni. A 278 vizsgálati alany válaszai alapján a politikus kedveltsége szignifikánsan csökkent az őt lejárató, hamis videó megnézése után – ám a politikus pártja iránti elkötelezettség változatlan maradt a kontrollállapothoz képest. A politikai színtér mellett a *deepfake*-kel előállított tartalmak a gazdaság befolyásolására is képesek (például egy tőzsdei hír elterjesztésével).

Az a képesség, hogy digitális szoftverek által készített képeket hozzunk létre, manipuláljuk és terjesszük őket, ma már mindenütt jelen van. Messaris a vizuális műveltség jelenségének két fontos következményét emeli ki.¹¹⁴ Egyrészt a digitális média nehezen felismerhető, árnyaltabb vizuális eszköztárat biztosít a képalkotók számára, másrészt azonban a digitális hálozatok sokkal könnyebbé tették a vizuális manipuláció felderítését, ebből fakadóan a vizuális műveltség is egyre inkább részét képezi a köztudatnak.

EREDMÉNYEK

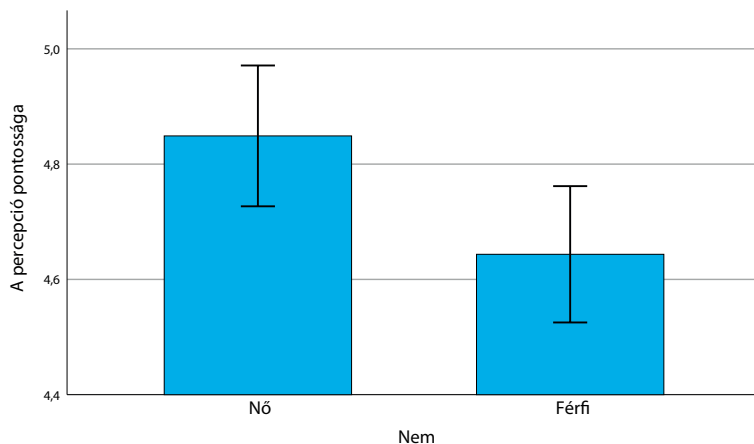
Kutatásunk záró szakaszában az elemzések fókuszában a valódi és MI-generált képek megkülönböztetésének képessége állt, amelyet a 8. fejezetben ismertetett percepciók mutatóval operacionalizáltunk. E feltáró elemzésekben

¹¹³ DOBBER et al. 2020.

¹¹⁴ MESSARIS 2012.

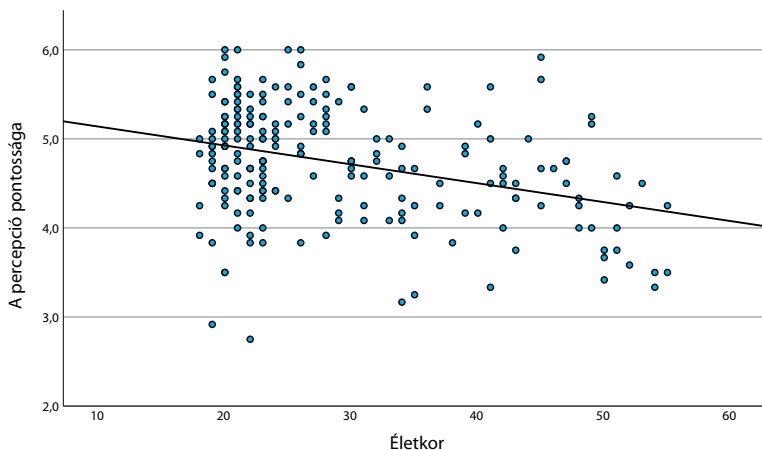
azt vizsgáltuk, mely egyéni jellemzők hozhatók összefüggésbe azzal, milyen hatékonyan képes a kitöltő e két ingertípust azonosítani és elkülöníteni egymástól. A védelemmotiváció négy összetevőjének (4. fejezet) és a személyiség öt fő dimenziójának (7. fejezet) egyike sem mutatott szignifikáns összefüggést a percepció pontosságával, ezért ebben a fejezetben a demográfiai jellemzők szerepének ismertetésére szorítkozunk.

Eredményeink szerint a női kitöltők szignifikánsan pontosabban ismerik fel az MI által generált képeket (átlag: 4,85), mint a férfiak (átlag: 4,64), bár a különbség kismértékű (Cohen-féle $d = 0,326$; $t[213] = 2,376$; $p = 0,018$). A 9.1. ábrán jól látható, hogy a 95%-os konfidenciaintervallumok csak kis mértékben fednek át egymással.



9.1. ábra: A nők pontosabban ítélik meg az embereket ábrázoló képek mesterséges vagy valódi eredetét (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés



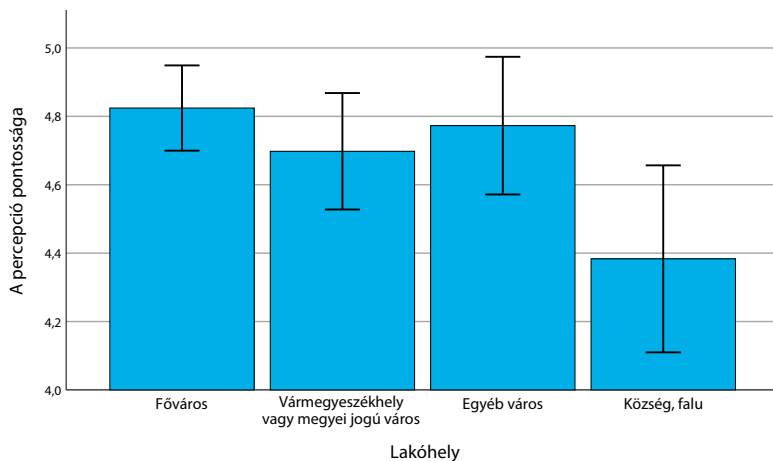
9.2. ábra: Az életkor korrelációja a percepció pontosságával

Forrás: saját szerkesztés

Hasonlóan szignifikáns összefüggést, közepes erősségű negatív korrelációt mutattunk ki az életkor és a percepció pontossága között ($r = -0,339$; $p < 0,001$). A 9.2. ábrán látható pontdiagramon jól kivehető a tendencia – amelyet a regressziós egyenes lejtése is megerősít –, hogy minél idősebb a felhasználó, annál alacsonyabb a percepció pontosság várható értéke. Az ábráról az is leolvasható, hogy az életkorban minden 10 évvel nagyjából 0,2 ponttal csökken az MI-detektálási képesség mutatójának várható értéke (a regressziós egyenes meredekségének pontos értéke: $b = -0,021219$). Míg a 20 évesek csoportjában a várható érték 4,93, az 50 évesek körében már csak 4,29.

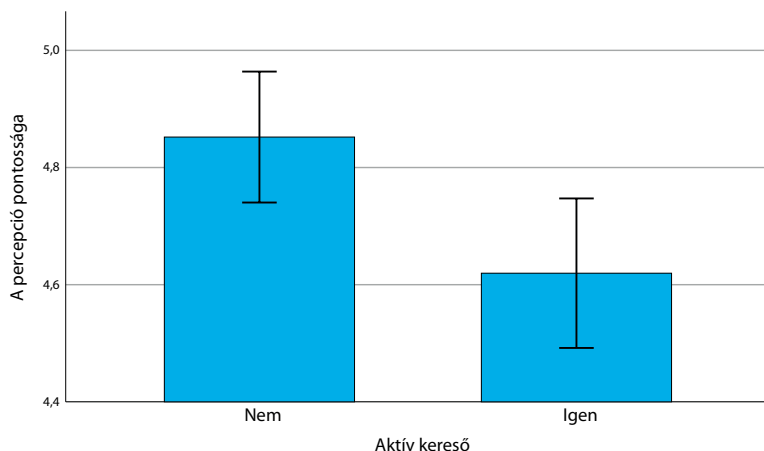
A nem és az életkoron túlmenően egy további fontos demográfiai változó, a lakóhely mérete is szignifikánsan összefügg a MI-generált képeken található árulkodó jelek felismerésének képességével. A varianciaanalízis szignifikáns eredményét követően ($F[3, 211] = 3,422$; $p = 0,018$) Bonferroni-korrektúrával *post hoc* elemzéseket végeztünk, amelyek eredménye azt mutatta, hogy a fővárosban élők MI-detektálási képessége (átlag: 4,82)

szignifikánsan jobb a falvakban, községekben élőkénél (átlag: 4,38). Ahogy a 9.3. ábrán is látszik, az átlagérték a vármegyeszékhelyeken (illetve egyéb megyei jogú városokban) és más városokban élők esetében a fővárosiak átlagos teljesítményéhez esik közelebb, az ő esetükben azonban a 95%-os konfidenciaintervallum már eléggé átfed a falvakban, községekben élőkével ahhoz, hogy a családonkénti hibaarány (*familywise error rate*) korrekciója után az átlagok közötti különbségek a $p < 0,05$ -ös szintén már nem szignifikánsak. Az adatsorunk tehát csupán arra szolgált bizonyítékot, hogy a fővárosiak és a falvakban, községekben lakók percepció pontossága között különbség mutatkozik. Ahhoz, hogy a másik két településtípussal kapcsolatban érdemi konklúziót vonhassunk le, nagyobb elemszámú mintára lenne szükségünk (ehhez valószínűleg elegendő lenne a falvakban, községekben élők elemszámát növelni, e demográfiai csoportot ugyanis a mintánkban mindössze 25 fő képviselte).



9.3. ábra: A percepció átlagos pontossága a lakóhely mérete szerint
(a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

Forrás: saját szerkesztés



9.4. ábra: Az aktív keresők kevésbé pontosan tudják megítélni egy kép valódiságát, mint azok, akik nem rendelkeznek rendszeres jövedelemmel (a hibahatárok 95%-os konfidenciaintervallumokat jelölnek)

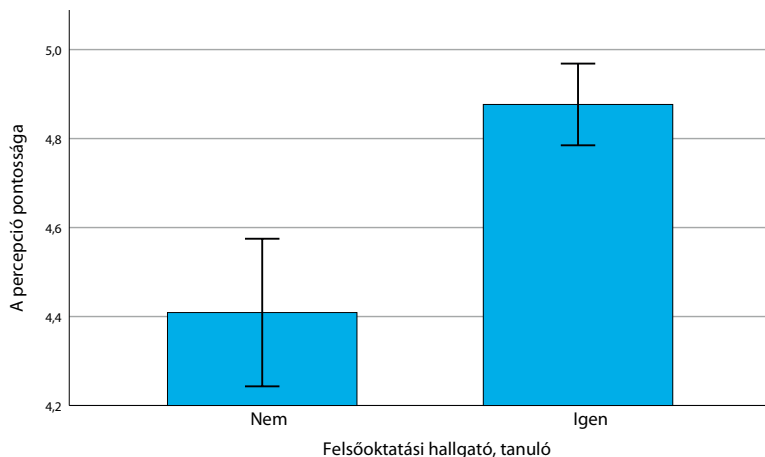
Forrás: saját szerkesztés

A független mintás t -próba eredményeiből az is látszik, hogy az aktív keresők vizuális MI-detektálási képessége (átlag: 4,62) szignifikánsan elmarad azokétól, akik (még) nem rendelkeznek rendszeres bérjövedelemmel (átlag: 4,85). A kismértékű (Cohen-féle $d = 0,37$), de erősen szignifikáns eltérést ($t[213] = 2,714, p = 0,007$) a 9.4. ábra szemlélteti.

Mielőtt azonban e részeredmény alapján elhamarkodottan arra a következtetésre jutnánk, hogy a rendszeres munka rontja az MI-ről árulkodó vizuális jegyek felismerésének képességét, fontos észrevennünk: az aktív kereső státusz az életkorral rendkívül erősen korrelál (a pont-biszeriális korrelációs együttható: $r_{pb} = 0,629; p < 0,001$). Ezért könnyen lehet, hogy az aktív kereső státusz hatása egyszerűen magyarázható azzal a körülménnyel, hogy a mintánkban az aktív keresők jellemzően az idősebb kitöltők voltak, és voltaképpen ebben az esetben is az életkor hatását látjuk érvényesülni. E magyarázatot egyértelműen megerősítik annak a többszörös lineáris regressziós modellnek az eredményei, amelyben a percepció pontosságát

az életkorból és az aktív kereső státuszból egyidejűleg próbáltuk meg bejósolni. Ilyen esetben az életkor szignifikáns magyarázó ereje megmarad (standardizált $\beta = -0,371$; $t[212] = -4,468$; $p < 0,001$) az aktív kereső státusz azonban ezen felül további magyarázó erővel nem rendelkezik (standardizált $\beta = 0,050$; $t[212] = ,608$; n. s.). Az elemzés alapján tehát megállapítható, hogy az aktív kereső státusz negatív hatása a percepció pontosságára kizárólag annak köszönhető, hogy az aktív keresők a mintánkban jellemzően idősebb személyek voltak, mint azok, akik rendszeres bérjövdelemmel nem rendelkeztek, és ezáltal az életkor hatása érvényesülhetett.

Az aktív kereső státuszhoz hasonló, de jóval nagyobb mértékű, szignifikáns különbség mutatkozott a mintában szereplő egyetemi hallgatók és a hallgatói státusszal nem rendelkezők között (Cohen-féle $d = -0,778$; $t[213] = -5,241$; $p < 0,001$). Ahogy a 9.5. ábra is mutatja, a felsőoktatási hallgatók MI-felismerési képessége szignifikánsan jobb (átlag: 4,88), mint azoké a kitöltőké, akik jelenleg nem egyetemi hallgatók (átlag: 4,41).



9.5. ábra: A felsőoktatási hallgatói státusszal rendelkezők pontosabban érzékelik a képek mesterséges vagy valódi eredetét

Forrás: saját szerkesztés

A fenti okfejtés logikája mentén gondolhatnánk, hogy e különbség magyarázatára is adódik az életkor mint harmadik változó hatása, különösen figyelembe véve, hogy a hallgatói státusz és az életkor közötti pont-biszeriális korreláció még erősebb, mint a fenti esetben ($r_{pb} = -0,669$; $p < 0,001$). Mégis, ha a fenti módszerrel, egy többszörös lineáris regressziós modellel vizsgáljuk az életkor és hallgatói státusz együttes hatását a percepció pontosságára, azt kapjuk, hogy mindkét tényező egymás mellett is szignifikáns marad, és nagyjából azonos mértékben járulnak hozzá a percepció képességben mutatózó egyéni különbségek magyarázatához (életkor: standardizált $\beta = -0,205$; $t[212] = -2,387$; $p = 0,018$; hallgatói státusz: standardizált $\beta = 0,201$; $t[212] = 2,344$; $p = 0,020$). Vagyis az életkor hatásának statisztikai kontrollja mellett azonos korcsoportba tartozó személyek esetében is elmondható, hogy a felsőoktatásban tanulók MI-felismerési képessége szignifikánsan jobb, mint azoké, akik nem felsőoktatási hallgatók.

A fentiekén kívül egyéb demográfiai változók nem mutattak összefüggést a percepció pontosságával. Nem mutatkoztak szignifikáns különbségek a különböző végzettségű (érettségi bizonyítvánnyal, alapképzésen, illetve mesterképzésen szerzett oklevéllel rendelkezők) között a tekintetben, milyen hatékonyan képesek megkülönböztetni az MI-generált képeket a valódi fényképektől.

Ehhez hasonlóan a kiberbiztonsági képzésen korábban részt vevő és ilyen tréningben nem részesült kitöltők között sem volt kimutatható különbség. Az előző fejezetek eredményeivel összevetve tehát az látszik, hogy míg a kitöltőink számára elérhető kiberbiztonsági képzések sikeresen alakítják ki azt a meggyőződést, hogy a résztvevők hatékonyan képesek védekezni az online visszaélések ellen (önhatékonyaság, 5. fejezet), valamint közvetlenül hozzájárulnak olyan védekező viselkedések kialakulásához, mint a rendszeres tájékozódás és az eszközök naprakészen tartása (3. fejezet), ezek a képzések önmagukban nem készítik fel a résztvevőket arra, hogy hatékonyabban legyenek képesek megkülönböztetni az MI által generált vizuális tartalmakat a valódi fényképektől.

ÖSSZEGZÉS

A valódi és a mesterséges intelligenciával generált, személyeket ábrázoló képek perceptuális megkülönböztetésének képessége eredményeink szerint öt változóval mutat szignifikáns összefüggést: (1) a nőkre jellemzőbb, mint a férfiakra, hogy a mesterséges képekre helyesen gyanakszanak és a valódi fényképekre nem; (2) a fiatalabbak nagyobb valószínűséggel azonosítják az MI által generált tartalmakat; (3) a fővárosban élők helyesebben ítélik meg a képek valódiságát, mint a falvakban, községekben lakók; (4) a rendszeres jövedelemmel nem rendelkezők percepciója pontosabb, mint az aktív keresőké, és (5) az egyetemi hallgatói státusszal rendelkezők hatékonyabban képesek megkülönböztetni a két ingertípust, mint akik (már) nem felsőoktatási hallgatók.

Az utóbbi két részeredmény között azonban fontos különbség, hogy míg az aktív kereső státusz hatása az életkor kontrollja mellett eltűnik – vagyis ezen összefüggés mögött valójában az életkor indirekt hatása húzódik meg –, az egyetemi hallgatók és a nem hallgatók közötti különbség az életkor kontrollja mellett is szignifikáns marad: az egyetemi hallgatói lét tehát még azonos életkorú egyének esetében is kedvezően befolyásolja a MI-generált képek helyes detektálásának képességét.

Kutatásunk válaszadóinak 43,7%-át (94 kitöltő) 18 és 25 év közötti, alapszakon tanuló egyetemi hallgatók teszik ki, ebből adódóan a kutatási minta kiemelkedő korcsoportja az 1995 és 2009 között született Z generáció (3. fejezet). Nem meglepő, hogy a Z generáció MI-felismerési képessége kiemelkedőbb, hiszen a fiatal közösségimédia-felhasználók köztudottan nyitottabbak a digitális újításokra, tudatosabbak a digitális eszközhasználatban, illetve sokkal intenzívebben van lehetőségük részt venni a digitális oktatás adta lehetőségekben. Emellett a „trendsetter”-nek is titulált Z generációt a kíváncsiság, a gyors alkalmazkodás és a fejlett digitális képesség jellemzi.¹¹⁵ Ezen jellemzők generációs különbségeket képeznek és megnehezítik a generációkon átívelő egységes digitális fejlődést.¹¹⁶ Ezt felismerve

¹¹⁵ OZIMEK et al. 2023; SZABÓ 2019.

¹¹⁶ KIRJAKOVSKI 2023.

az OECD által kezdeményezett *Az oktatás és készségek jövője 2030* arra fekteti a hangsúlyt, hogy támogassa az országokat oktatási rendszereik átalakításában, figyelembe véve azon kiemelten fontos digitális kompetenciákat, amelyekre a diákoknak és tanároknak egyaránt szükségük van a hatékony és biztonságos médiahasználathoz.¹¹⁷

A kapott eredmény, miszerint a nők nagyobb arányban voltak képesek kiszűrni a manipulált tartalmakat, összhangban áll a Veszelszki–Horváth–Kovács kutatócsoport korábbi kimutatásaival, amelyben mesterségesintelligencia-alapú *deepfake* videókat vizsgáltak.¹¹⁸ A kutatás célja az volt, hogy kiderítsék, a befogadók képesek-e felismerni a manipulált képeket, és ha igen, meg tudják-e határozni a manipuláció pontos helyét a bemutatott vizuális tartalmakon. Online kérdőív segítségével kiemelkedően nagy mintán (N = 10.380) vizsgálták a kép- és videomanipulált tartalmak befogadóra tett hatását. A minta 9 videóból és 9 portréfotóból állt. A kutatás egyik fő eredménye, hogy a videomanipuláció felismerése kevésbé okoz nehézséget a befogadónak, mint a képmanipuláció kiszűrése. Mivel az arc minősül az egyik legerősebb vizuális ingernek, így nem meglepő, hogy az alanyok 60%-a számolt be arról, hogy a szem, hajszín és arcszimmetria nem tűnt számukra természetesnek a bemutatott fotókon.

A közösségi felületek egyik előnye, hogy személyre szabhatóak, ezáltal minden felhasználó saját maga formálhatja az online megjelenített énképét. Ez azonban hátrányként is értelmezhető, hiszen a felhasználók a legkedvezőbb oldalunkat szeretnék megmutatni, ezáltal a képmanipulálás módszereihez folyamodnak, ami bumerángszerűen visszahathat önértékelésünkre és torzíthatja azt.¹¹⁹ 2024 novemberében a TikTok megpróbálta megtiltani a felhasználóinak, hogy szépségfiltereket alkalmazzanak képeiken (nem sikerült ezt a célt elérni). A közösségi felület azzal érvelt, hogy az új szabályozással próbálják enyhíteni az irreális elvárásokat, amelyeket a közösségi média által vezérelt szépségipar közvetít a fiatal lányok felé.¹²⁰

¹¹⁷ OECD 2024.

¹¹⁸ VESZELSZKI–HORVÁTH–KOVÁCS 2022.

¹¹⁹ OZIMEK et al. 2023.

¹²⁰ IJAZ 2024.

Mills és társainak vizsgálata is azt támasztja alá, hogy a női felhasználók sokkal nagyobb arányban szerkesztik saját magukról készített képeiket, mint férfi társaik, sőt, önmagában egy szelfi közzététele is komoly szorongást és egyéb pszichológiai hatásokat válthat ki a nők esetében.¹²¹ Egy hasonló kutatásban arra is fény derült, hogy a nők közösségi médiában közzétett képei pozitívan korrelálnak a saját testükről alkotott véleménynel.¹²²

Masahiro Mori 55 évvel ezelőtt, 1970-ben megalkotta a „hátborzongató völgy” (*uncanny valley*) hipotézist, amely arra épül, hogy minél jobban hasonlítanak a robotok az emberekre, annál erősebben képesek rokonszenvet kiváltani belőlünk. Ám ha egy ponton túlságosan hasonlók lesznek az emberekhez, akkor elveszítjük feléjük a szimpátiát, és egy „hátborzongató völgy”-be kerülünk.¹²³ Bár sosem támasztották alá, mégis talán most járunk a legközelebb a hipotézis igazolásához.

A digitális képalkotás korszakába lépve nem csupán a fotós és videós szerkesztői gyakorlatokban történtek radikális változások, hanem ezáltal egy új típusú társadalmi kétely is megjelent (a kétely kételye paradoxon).¹²⁴ Amíg a *seeing is believing* mondás évtizedeken keresztül meghatározta a vizuális kommunikáció hitelességét, addig a közösségimédia-platformokon vírusszerűen elterjedt MI-alapú képmanipuláció egyre inkább megkérdőjelezi a vizuális produktumok valóságát.

A fejezet arra a kérdésre kereste a választ, milyen egyéni jellemzők hozhatók összefüggésbe azzal, hogy valaki hatékonyabban vagy kevésbé hatékonyan képes MI által létrehozott vizuális tartalmakat megkülönböztetni a valódi fényképektől. Eredményeinkből arra következtethetünk, hogy nemcsak a vizuális tartalmakban fellelhető jellemzők befolyásolhatják a felismerőképességet, hanem az egyén személyes, tanulmányi háttére is. A kutatásból kiderült, hogy a 18–25 év közötti Z generációs hallgatók voltak leginkább képesek felismerni a képmanipulációt, közülük is a nők teljesítettek jobban. Azt is megtudtuk, hogy azon kitöltők, akik aktívan

¹²¹ MILLS et al. 2018.

¹²² MCGOVERN–COLLINS–DUNNE 2022.

¹²³ WANG–LILIENFELD–ROCHAT 2015.

¹²⁴ CHESNEY–CITRON 2018; VESZELSZKI 2023.

a felsőoktatásban tanulnak, hatékonyabban tudják megkülönböztetni az MI-generált képeket a valódi fotóktól.

„Ha az emberek valótlanok, minden más is valótlanra válik” – fogalmazza meg sikerkönyvében Jaron Lanier számítógépfilozófia-író, futurista.¹²⁵ A vizuális percepció kiszélesedése ma már a digitális-vizuális műveltség elsajátítását, fejlesztését is szorgalmazza, hiszen az ezzel kapott tudás elősegíti a közösségi platformok aktív felhasználóinak tudatos viselkedését, etikus képhasználatát, valamint a valós és valótlan tartalmak hatékony megkülönböztetését.

¹²⁵ LANIER 2020: 54.

FELHASZNÁLT IRODALOM

- AKTER, Shahriar et al. (2022): Reconceptualizing Cybersecurity Awareness Capability in the Data-Driven Digital Economy. *Annals of Operations Research*. Online: <https://doi.org/10.1007/s10479-022-04844-8>
- AL ABDULWAHID, Abdulwahid et al. (2015): Security, Privacy and Usability – A Survey of Users’ Perceptions and Attitudes. In FISCHER-HÜBNER, Simone – LAMBRINOUDAKIS, Costas – LÓPEZ, Javier (szerk.): *Trust, Privacy and Security in Digital Business. TrustBus 2015. Lecture Notes in Computer Science*, vol 9264. Cham: Springer. Online: https://doi.org/10.1007/978-3-319-22906-5_12
- ALHAMAD, Hamza – DONYAI, Parastou (2021): The Validity of the Theory of Planned Behaviour for Understanding People’s Beliefs and Intentions toward Reusing Medicines. *Pharmacy*, 9(1), 58. Online: <https://doi.org/10.3390/pharmacy9010058>
- ALMANSRI, Afrah – AL-EMRAN, Mostafa – SHAALAN, Khaled (2023): Exploring the Frontiers of Cybersecurity Behavior: A Systematic Review of Studies and Theories. *Applied Sciences*, 13(9), 5700. Online: <https://doi.org/10.3390/app13095700>
- ARON, Jacob (2012): Passwords for Over-55s Are Twice as Secure as Those Chosen by Teenagers. *New Scientist*, 214(2866), 22. Online: [https://doi.org/10.1016/S0262-4079\(12\)61347-5](https://doi.org/10.1016/S0262-4079(12)61347-5)
- BAIG, Anas (2024): 25 Essential Cybersecurity Tips and Best Practices for Your Business. *LevelBlue*, 2024. március 13. Online: <https://levelblue.com/blogs/security-essential-s/25-essential-cybersecurity-tips-and-best-practices-for-your-business>
- BARON, Reuben M. – KENNY, David A. (1986): The Moderator–Mediator Variable Distinction in Social Psychological Research: Conceptual, Strategic, and Statistical Considerations. *Journal of Personality and Social Psychology*, 51(6), 1173–1182. Online: <https://doi.org/10.1037/0022-3514.51.6.1173>
- BAUER, Johannes M. (2018): The Internet and Income Inequality: Socio-Economic Challenges in a Hyperconnected Society. *Telecommunications Policy*, 42(4), 333–343. Online: <https://doi.org/10.1016/j.telpol.2017.05.009>
- BERGER, Arthur Asa (2012): *Seeing Is Believing. An Introduction To Visual Communication*. New York: McGraw-Hill Education.
- BOEHMER, Jan et al. (2015): Determinants of Online Safety Behaviour: Towards an Intervention Strategy for College Students. *Behaviour & Information Technology*, 34(10), 1022–1035. Online: <https://doi.org/10.1080/0144929X.2015.1028448>
- CARLINI, Nicholas et al. (2020): Extracting Training Data from Large Language Models. *USENIX Security Symposium*. Online: <https://www.usenix.org/conference/usenix-security21/presentation/carlini-extracting>

- CASAS, Andreu – WILLIAMS, Nora Webb (2018): Images that Matter: Online Protests and the Mobilizing Role of Pictures. *Political Research Quarterly*, 72(2), 360–375. Online: <https://doi.org/10.1177/1065912918786805>
- CHAUDHARY, Sunil (2024): Driving Behaviour Change with Cybersecurity Awareness. *Computers & Security*, 142(C), 103858. Online: <https://doi.org/10.1016/j.cose.2024.103858>
- CHESNEY, Robert – CITRON, Danielle Keats (2019): Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review* 107(6), 1753–1820. Online: <https://www.jstor.org/stable/26891938?seq=1>
- CHY, Md Kamrul Hasan (2024): Proactive Fraud Defense: Machine Learning's Evolving Role in Protecting Against Online Fraud. Online: <https://doi.org/10.48550/arXiv.2410.20281>
- CINELLI, Matteo et al. (2020): The COVID-19 Social Media Infodemic. *Scientific Reports*, 10(1), 16598, Online: <https://doi.org/10.1038/s41598-020-73510-5>
- CISMARU, Magdalena et al. (2011): “Act on Climate Change”: An Application of Protection Motivation Theory. *Social Marketing Quarterly*, 17(3), 62–84. Online: <https://doi.org/10.1080/15245004.2011.595539>
- Coursera Staff (2024): 9 Cybersecurity Best Practices for Businesses in 2024. *Coursera*, 2024. november 10. Online: <https://www.coursera.org/articles/cybersecurity-best-practices>
- CROSS, Cassandra (2022): Using Artificial Intelligence (AI) and Deepfakes to Deceive Victims: The Need to Rethink Current Romance Fraud Prevention Messaging. *Crime Prev Community Saf*, 24(1), 30–41. Online: <https://doi.org/10.1057/s41300-021-00134-w>
- DAWSON, Jessica – THOMSON, Robert (2018): The Future Cybersecurity Workforce: Going Beyond Technical Skills for Successful Cyber Performance. *Frontiers in Psychology*, 9. Online: <https://doi.org/10.3389/fpsyg.2018.00744>
- DEBES, John L. (1969): The Loom of Visual Literacy: An Overview. *Audiovisual Instruction*, 14(8), 25–27.
- DHARIWAL, Prafulla – NICHOL, Alex (2024): Diffusion Models Beat GANs on Image Synthesis. *Proceedings of the 35th International Conference on Neural Information Processing Systems*. Red Hk, NY, USA: Curran Associates Inc., 8780–8794. Online: <https://doi.org/10.48550/arXiv.2105.05233>
- DOBBER, Tom – METOUI, Nadia – TRILLING, Damian – HELBERGER, Natali – DE VREESE, Claes 2020: Do (Microtargeted) Deepfakes Have Real Effects on Political Attitudes? *The International Journal of Press/Politics*, 26(1), 1–23. Online: <https://doi.org/10.1177/1940161220944364>
- DODEL, Matias – MESCH, Gustavo (2019): An Integrated Model for Assessing Cyber-Safety Behaviors: How Cognitive, Socioeconomic, and Digital Determinants Affect Diverse Safety Practices. *Computers & Security*, 86, 75–91. Online: <https://doi.org/10.1016/j.cose.2019.05.023>
- DROGKARIS, Prokopios – BOURKA, Athena szerk. (2018): *Cybersecurity Culture Guidelines: Behavioural Aspects of Cybersecurity*. Marousi: European Union Agency For Network and Information Security. Online: <https://data.europa.eu/doi/10.2824/324042>

- ESZTERI Dániel (2020): Elosztott mesterséges intelligencia fejlesztés blokklánc alapon az adatvédelem érvényesülése érdekében. *Pro Futuro*, 10(1), 9–27. Online: <https://doi.org/10.26521/Profuturo/2020/1/7554>
- FERRARA, Emilio (2017): Disinformation and Social Bot Operations in the Run up to the 2017 French Presidential Election. *First Monday*, 22(8). Online: <https://doi.org/10.5210/fm.v22i8.8005>
- FIELD, Andy P. (2009): *Discovering Statistics Using SPSS (and Sex and Drugs and Rock 'n' Roll)*. London: Sage Publications Ltd.
- FRANK, Joel et al. (2024): A Representative Study on Human Detection of Artificially Generated Media Across Countries. 2024 *IEEE Symposium on Security and Privacy (SP)*, 55–73. Online: <https://doi.org/10.48550/arXiv.2312.05976>
- GALLESE, Vittorio (2024): Digital Visions: The Experience of Self and Others in the Age of the Digital Revolution. *International Review of Psychiatry*, 36(6), 656–666. Online: <https://doi.org/10.1080/09540261.2024.2355281>
- GOODFELLOW, Ian et al. (2020): Generative Adversarial Networks. *Communications of the ACM*, 63(11), 139–144. Online: <https://doi.org/10.1145/3422622>
- GRUBER, Christiane – HAUGBOLLE, Sune szerk. (2013): *Visual Culture in the Modern Middle East: Rhetoric of the Image*. Bloomington: Indiana University Press.
- HAMELEERS, Michael – VAN DER MEER, Toni G. L. A. – DOBBER, Tom (2024): They Would Never Say Anything Like This! Reasons To Doubt Political Deep-fakes. *European Journal of Communications*, 39(1), 56–70. Online: <https://doi.org/10.1177/02673231231184703>
- HAMOUD, Aymen – AÏMEUR, Esma (2020): Handling User-Oriented Cyber-Attacks: STRIM, a User-Based Security Training Model. *Frontiers of Computer Science*, 2. Online: <https://doi.org/10.3389/fcomp.2020.00025>
- HANCOCK, Jeffrey T. – NAAMAN, Mor – LEVY, Karen (2020): AI-Mediated Communication: Definition, Research Agenda, and Ethical Considerations. *Journal of Computer-Mediated Communication*, 25(1), 89–100. Online: <https://doi.org/10.1093/jcmc/zmz022>
- HANUS, Bartłomiej – WU, Yu “Andy” (2016): Impact of Users’ Security Awareness on Desktop Security Behavior: A Protection Motivation Theory Perspective. *Information Systems Management*, 33(1), 2–16. Online: <https://doi.org/10.1080/10580530.2015.1117842>
- HARRIS, Christine – JENKINS, Michael (2006): Gender Differences in Risk Assessment: Why do Women Take Fewer Risks than Men? *Judgment and Decision Making*, 1(1), 48–63. Online: <https://doi.org/10.1017/S1930297500000346>
- HEFFERNAN, Michael (2015): The Interrogation of Sándor Radó: Geography, Communism and Espionage Between World War Two and the Cold War. *Journal of Historical Geography*, 47, 74–88. Online: <https://doi.org/10.1016/j.jhbg.2014.12.004>
- HULL, James L. (2017): *Analyst Burnout in the Cyber Security Operation Center-CSOC: A Phenomenological Study*. PhD-dissertáció. Colorado Technical University.

- HUMAIDI, Norshima – ABDALLAH ALGHAZO, Saif Hussein (2022): Procedural Information Security Countermeasure Awareness and Cybersecurity Protection Motivation in Enhancing Employee's Cybersecurity Protective Behaviour. 2022 10th International Symposium on Digital Forensics and Security (ISDFS), 1–10. Online: <https://doi.org/10.1109/ISDFS55398.2022.9800834>
- HUXLEY, Aldous (1974): *The Art of Seeing*. London: Chatto & Windus.
- IJAZ, Ezza (2024): TikTok Will Ban Beauty Filters for Users Under 18 to Address the Harmful Impact of Unrealistic Beauty Standards on Self-Esteem. *WCCFTECH*, 2024. november 28. Online: <https://wccftech.com/tiktok-will-ban-beauty-filters-for-users-under-18/>
- Institute of Data (2023): How to Stay Up to Date with Cyber Security. *Institute of Data*, 2023. október 9. Online: <https://www.institutedata.com/blog/how-to-stay-up-to-date-with-cyber-security/>
- JEONG, Jongkil Jay et al. (2019): Towards an Improved Understanding of Human Factors in Cybersecurity. 2019 IEEE 5th International Conference on Collaboration and Internet Computing (CIC), 338–345. Online: <https://doi.org/10.1109/CIC48465.2019.00047>
- JOHNSON, John A. (2014): Measuring Thirty Facets of the Five Factor Model with a 120-Item Public Domain Inventory: Development of the IPIP-Neo-120. *Journal of Research in Personality*, 51, 78–89. Online: <https://doi.org/10.1016/j.jrp.2014.05.003>
- KEITH, Timothy Z. (2019): *Multiple Regression and Beyond. An Introduction to Multiple Regression and Structural Equation Modeling*. New York: Routledge.
- KHADER, Mohammed – KARAM, Marcel – FARES, Hanna (2021): Cybersecurity Awareness Framework for Academia. *Information*, 12(10), 417. Online: <https://doi.org/10.3390/info12100417>
- KHAN, Naurin Farooq – IKRAM, Naveed – SALEEM, Sumera (2024): Effects of Socioeconomic and Digital Inequalities on Cybersecurity in a Developing Country. *Security Journal*, 37, 214–244. Online: <https://doi.org/10.1057/s41284-023-00375-4>
- KIETZMANN, Jan et al. (2019): Deepfakes: Trick or Treat? *Business Horizons*, 63(2), 135–146. Online: <https://doi.org/10.1016/j.bushor.2019.11.006>
- KIRAN, Uzma et al. (2025): Explanatory and Predictive Modeling of Cybersecurity Behaviors Using Protection Motivation Theory. *Computers & Security*, 149, 104204. Online: <https://doi.org/10.1016/j.cose.2024.104204>
- KIRJAKOVSKI, Atanas (2023): Rethinking Perception and Cognition in the Digital Environment. *Frontiers in Cognition*, 2. Online: <https://doi.org/10.3389/fcogn.2023.1266404>
- KIRSCHNER, Julianna (2024): The Newest Generational Divide: Social Media Influence on Digital Natives. *Visual Communication Quarterly*, 31(3), 199–205. Online: <https://doi.org/10.1080/15551393.2024.2396791>
- LADEVIĆ, Adrian Lokner et al. (2024): Detection of AI-Generated Synthetic Images with a Lightweight CNN. *AI*, 5(3), 1575–1593. Online: <https://doi.org/10.3390/ai5030076>
- LANIER, Jaron (2020): *Miért töröld magad azonnal a közösségi oldalakról? 10 érv*. Budapest: Európa.

- LAZARD, Allison J. – BOCK, Mary A. – MACKERT, Michael S. (2020): Impact of Photo Manipulation and Visual Literacy on Consumers' Responses to Persuasive Communication. *Journal of Visual Literacy*, 39(2), 90–110. Online: <https://doi.org/10.1080/1051144X.2020.1737907>
- LEBEK, Benedikt et al. (2014): Information Security Awareness and Behavior: A Theory-Based Literature Review. *Management Research Review*, 37(12), 1049–1092. Online: <https://doi.org/10.1108/MRR-04-2013-0085>
- LI, Ling et al. (2019): Investigating the Impact of Cybersecurity Policy Awareness on Employees' Cybersecurity Behavior. *International Journal of Information Management*, 45, 13–24. Online: <https://doi.org/10.1016/j.ijinfomgt.2018.10.017>
- LU, Zeyu et al. (2023): Seeing Is Not Always Believing: Benchmarking Human and Model Perception of AI-Generated Images. *Proceedings of the 37th International Conference on Neural Information Processing Systems*, 25435–25447. Online: <https://doi.org/10.48550/arXiv.2304.13023>
- MAPLES-KELLER, Jessica L. et al. (2019): Using Item Response Theory to Develop a 60-Item Representation of the NEO PI-R Using the International Personality Item Pool: Development of the IPIP-NEO-60. *Journal of Personality Assessment*, 101(1), 4–15. Online: <https://doi.org/10.1080/00223891.2017.1381968>
- MARIANO, João et al. (2021): Too Old for Technology? Stereotype Threat and Technology Use by Older Adults. *Behaviour & Information Technology*, 41(7), 1503–1514. Online: <https://doi.org/10.1080/0144929X.2021.1882577>
- MARIKYAN, Davit – PAPAGIANNIDIS, Savvas (2023): Protection Motivation Theory: A review. In PAPAGIANNIDIS, Savvas (szerk.): *TheoryHub Book*, 78–93. Online: <https://open.ncl.ac.uk/theory-library/TheoryHubBook.pdf>
- MARTINEAU, Melissa – SPIRIDON, Elena – AIKEN, Mary (2023): A Comprehensive Framework for Cyber Behavioral Analysis Based on a Systematic Review of Cyber Profiling Literature. *Forensic Sciences*, 3(3), 452–477. Online: <https://doi.org/10.3390/forensicsci3030032>
- MARTIN-RODRIGUEZ, Fernando – GARCIA-MOJON, Rocio – FERNANDEZ-BARCIELA, Monica (2023): Detection of AI-Created Images Using Pixel-Wise Feature Extraction and Convolutional Neural Networks. *Sensors*, 23(22), 9037. Online: <https://doi.org/10.3390/s23229037>
- MARWICK, Alice E. – MILLER, Ross (2014): *Online Harassment, Defamation, and Hateful Speech: A Primer of the Legal Landscape*. Fordham Center on Law and Information Policy Report No. 2. Online: <https://ssrn.com/abstract=2447904>
- MCCRAE, Robert R. – COSTA, Paul T. (1997): Personality Trait Structure as a Human Universal. *American Psychologist*, 52(5), 509–516. Online: <https://doi.org/10.1037/0003-066X.52.5.509>
- MC ELROY, James C. et al. (2007): Dispositional Factors in Internet Use: Personality versus Cognitive Style. *MIS Quarterly*, 31(4), 809–820. Online: <https://doi.org/10.2307/25148821>

- MCGOVERN, Orla – COLLINS, Rebecca – DUNNE, Simon (2022): The Associations Between Photo-Editing and Body Concerns Among Females: A Systematic Review. *Body Image*, 43, 504–517. Online: <https://doi.org/10.1016/j.bodyim.2022.10.013>
- MENARD, Philip – BOTT, Gregory J. – CROSSLER, Robert E. (2017): User Motivations in Protecting Information Security: Protection Motivation Theory Versus Self-Determination Theory. *Journal of Management Information Systems*, 34 (4), 1203–1230. Online: <https://doi.org/10.1080/07421222.2017.1394083>
- MESSARIS, Paul (2012): Visual “Literacy” in the Digital Age. *Review of Communication*, 12(2), 101–117. Online: <https://doi.org/10.1080/15358593.2011.653508>
- MIJWIL, Maad – SALEM, Israa Ezzat – ISMAEEL, Marwa M. (2023): The Significance of Machine Learning and Deep Learning Techniques in Cybersecurity: A Comprehensive Review. *Iraqi Journal for Computer Science and Mathematics*, 4(1), 87–101. Online: <https://doi.org/10.52866/ijcsm.2023.01.01.008>
- MILLS, Jennifer S. et al. (2018): “Selfie” Harm: Effects on Mood and Body Image in Young Women. *Body Image*, 27, 86–92. Online: <https://doi.org/10.1016/j.bodyim.2018.08.007>
- MIRSKY, Yisroel – LEE, Wenke (2021): The Creation and Detection of Deepfakes: A Survey. *ACM Computing Surveys*, 54(1), 7:1–7:41. Online: <https://doi.org/10.1145/3425780>
- MOU, Jian et al. (2022): A Test of Protection Motivation Theory in the Information Security Literature: A Meta-Analytic Structural Equation Modeling Approach. *Journal of the Association for Information Systems*, 23(1), 196–236. Online: <https://doi.org/10.17705/1jais.00723>
- MOUSTAFA, Ahmed A. – BELLO, Abubakar – MAURUSHAT, Alana (2021): The Role of User Behaviour in Improving Cyber Security Management. *Frontiers in Psychology*, 12. Online: <https://doi.org/10.3389/fpsyg.2021.561011>
- NEWMAN, Eryn J. et al. (2012): Nonprobative Photographs (Or Words) Inflate Truthiness. *Psychonomic Bulletin & Review*, 19(5), 969–974. Online: <https://doi.org/10.3758/s13423-012-0292-0>
- NEWTON, Julianne H. (2006): Influences of Digital Imaging on the Concept of Photographic Truth. In MESSARIS, Paul – HUMPHREYS, Lee (szerk.): *Digital Media: Transformations in Human Communication*. New York: Peter Lang, 3–14.
- NIGHTINGALE, Sophie J. – FARID, Hany (2022): AI-Synthesized Faces Are Indistinguishable from Real Faces and More Trustworthy. *Proceedings of the National Academy of Sciences*, 119(8), e2120481119. Online: <https://doi.org/10.1073/pnas.2120481119>
- NIGHTINGALE, Sophie J. – WADE, Kimberley A. – WATSON, Derrick G. (2017): Can People Identify Original and Manipulated Photos of Real-World Scenes? *Cognitive Research: Principles and Implications*, 2, 30. Online: <https://doi.org/10.1186/s41235-017-0067-2>
- NMHH (2024): Illúzió vagy valóság? – A deepfake-dilemma. nmhh.hu, 2024. augusztus 9. Online: https://nmhh.hu/cikk/248038/Illuzio_vagy_valosag__A_deepfakedilemma

- NOBLES, Calvin (2021): The Human Factors Series: Burnout and Fatigue are Sustained Problems in Cybersecurity. *LinkedIn*, 2021. február 8. Online: <https://www.linkedin.com/pulse/human-factors-series-burnout-fatigue-sustained-calvin-nobles-ph-d/>
- NOBLES, Calvin (2022): Stress, Burnout, and Security Fatigue in Cybersecurity: A Human Factors Problem. *HOLISTICA – Journal of Business and Public Administration*, 13(1), 49–72. Online: <https://doi.org/10.2478/hjbpa-2022-0003>
- NURSE, Jason R. C. (2019): Cybercrime and You: How Criminals Attack and the Human Factors That They Seek to Exploit. In ATTRILL-SMITH, Alison et al. (szerk.): *The Oxford Handbook of Cyberpsychology*. Oxford: Oxford University Press, 663–690.
- OECD (2024): *Future of Education and Skills 2030*. Online: <https://www.oecd.org/en/about/projects/future-of-education-and-skills-2030.html>
- OZIMEK, Phillip et al. (2023): How Photo Editing in Social Media Shapes Self-Perceived Attractiveness and Self-Esteem via Self-Objectification and Physical Appearance Comparisons. *BMC Psychology*, 11, 99. Online: <https://doi.org/10.1186/s40359-023-01143-0>
- PALETZ, Susannah B. F. et al. (2023): Emotional Content and Sharing on Facebook: A Theory Cage Match. *Science Advances*, 9(39), eade9231. Online: <https://doi.org/10.1126/sciadv.ade9231>
- PARRY, Zachariah B. (2009): Digital Manipulation and Photographic Evidence: Defrauding the Courts One Thousand Words at a Time. *Journal of Law, Technology & Policy*, 2009(1), 175.
- POWELL, Thomas E. et al. (2015): A Clearer Picture: The Contribution of Visuals and Text to Framing Effects. *Journal of Communication*, 65(6), 997–1017. Online: <https://doi.org/10.1111/jcom.12184>
- RICKER, Jonas et al. (2024): AI-Generated Faces in the Real World: A Large-Scale Case Study of Twitter Profile Images. In *Proceedings of the 27th International Symposium on Research in Attacks, Intrusions and Defenses (RAID '24)*. New York: Association for Computing Machinery, 513–530. Online: <https://doi.org/10.1145/3678890.3678922>
- ROGERS, Ronald W. (1975): Protection Motivation Theory of Fear Appeals and Attitude Change. *Journal of Psychology: Interdisciplinary and Applied*, 91(1), 93–114. Online: <https://doi.org/10.1080/00223980.1975.9915803>
- ROHAN, Rohani et al. (2021): Understanding of Human Factors in Cybersecurity: A Systematic Literature Review. *2021 International Conference on Computational Performance Evaluation (ComPE)*, 133–140. Online: <https://doi.org/10.1109/ComPE53109.2021.9752358>
- ROMBACH, Robin et al. (2022): High-Resolution Image Synthesis with Latent Diffusion Models. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 10684–10695. Online: <https://doi.org/10.48550/arXiv.2112.10752>
- SHAUKAT, Kamran et al. (2020): A Survey on Machine Learning Techniques for Cyber Security in the Last Decade. *IEEE Access*, 8, 222310–222354. Online: <https://doi.org/10.1109/ACCESS.2020.3041951>

- SNIDER, Keren L. G. et al. (2021): Cyberattacks, Cyber Threats, and Attitudes Toward Cybersecurity Policies. *Journal of Cybersecurity*, 7(1), tyabo19. Online: <https://doi.org/10.1093/cybsec/tyabo19>
- SPALTER, Anne M. – VAN DAM, Andries (2008): Digital Visual Literacy. *Theory Into Practice*, 47(2), 93–101. Online: <https://doi.org/10.1080/00405840801992256>
- SUTTON, Stephen (2001): Health Behavior: Psychosocial Theories. In SMELSER, Neil J. – BALTES, Paul B. (szerk.): *International Encyclopedia of the Social & Behavioral Sciences*. Oxford: Elsevier, 6499–6506. Online: <https://doi.org/10.1016/Bo-08-043076-7/03872-9>
- ŠVÁBENSKÝ, Valdemar et al. (2021): Cybersecurity Knowledge and Skills Taught in Capture the Flag Challenges. *Computers & Security*, 102, 102154. Online: <https://doi.org/10.1016/j.cose.2020.102154>
- SZABÓ, Csilla Marianna (2019): Digital Competence of Teachers – How Do We Teach Generation Z? In ANDRÁS, István – RAJCSÁNYI-MOLNÁR, Mónika (szerk.): *East West Cohesion III. Strategic Study Volumes*. Subotica: Čikoš Group, 197–206
- SZOKOLSZKY Ágnes (2004): *Kutatómunka a pszichológiában. Metodológia, módszerek, gyakorlat*. Budapest: Osiris.
- TeamPassword (2024): 10 Cybersecurity Best Practices Your Employees Must Follow. *TeamPassword*, 2024. szeptember 17. Online: <https://teampassword.com/blog/cyber-security-best-practices-for-employees>
- THOMAS, Brian (2024): 5 Shocking IT & Cybersecurity Burnout Statistics. *Bitsight*, 2024. január 11. Online: <https://www.bitsight.com/blog/5-shocking-it-cybersecurity-burnout-statistics>
- VANNUCCI, Anna et al. (2020): Social Media Use and Risky Behaviors in Adolescents: A Meta-Analysis. *Journal of Adolescence*, 79(1), 258–274. Online: <https://doi.org/10.1016/j.adolescence.2020.01.014>
- VESZELSZKI, Ágnes – HORVÁTH, Evelin – KOVÁCS, Gábor (2022): New Media Literacy in Light of Image Manipulation and Deepfake Technology. In MAMMADOV, Azad – LEWANDOWSKA-TOMASZCZYK, Barbara (szerk.): *Analysing Media Discourse: Traditional and New*. Newcastle upon Tyne: Cambridge Scholars Publishing, 148–178.
- VESZELSZKI Ágnes (2021): deepFAKEnews: Az információmanipuláció új módszerei. In BALÁZS László (szerk.): *Digitális kommunikáció és tudatosság*. Budapest: Hungarovox, 93–105.
- VESZELSZKI Ágnes (2023): Deepfake: kételkedés a kételyben. In ACZÉL Petra – VESZELSZKI Ágnes (szerk.): *Deepfake: a valótlán valóság*. Budapest: Gondolat, 13–31.
- VOSOUGHI, Soroush – ROY, Deb – ARAL, Sinan (2018): The Spread of True and False News Online. *Science*, 359 (6380), 1146–1151. Online: <https://doi.org/10.1126/science.aap9559>
- WALKER, Sidney – WALKER, Sydney (2004): Artmaking in an Age of Visual Culture: Vision and Visuality. *Visual Arts Research*, 30(2), 23–37. Online: <http://www.jstor.org/stable/20715350>

- WANG, Shensheng – LILIENFELD, Scott O. – ROCHAT, Philippe (2015): The Uncanny Valley: Existence and Explanations. *Review of General Psychology*, 19(4), 393–407. Online: <https://doi.org/10.1037/gpr0000056>
- WANG, Sheng-Yu et al. (2020): CNN-Generated Images Are Surprisingly Easy to Spot... for Now. 2020 *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 8695–8704. Online: <https://doi.org/10.48550/arXiv.1912.11035>
- WANG, Wentao et al. (2023): DeepArt: A Benchmark to Advance Fidelity Research in AI-Generated Content. Online: <https://doi.org/10.48550/arXiv.2312.10407>
- WOLFF, Josephine (2021): How Is Technology Changing the World, and How Should the World Change Technology? *Global Perspectives*, 2(1), 27353. Online: <https://doi.org/10.1525/gp.2021.27353>
- YOUSSEF, Adel Ben – DAHMANI, Mounir – RAGNI, Ludovic (2022): ICT Use, Digital Skills and Students' Academic Performance: Exploring the Digital Divide. *Information*, 13(3), 129. Online: <https://doi.org/10.3390/info13030129>
- YU, Joanne et al. (2024): Artificial Intelligence-Generated Virtual Influencer: Examining the Effects of Emotional Display on User Engagement. *Journal of Retailing and Consumer Services*, 76, 103560. Online: <https://doi.org/10.1016/j.jretconser.2023.103560>
- YU, Yanqiu – LAU, Mason M. C. – LAU, Joseph T. F. (2022): Application of the Protection Motivation Theory to Understand Determinants of Compliance with the Measure of Banning Gathering Size >4 in All Public Areas for Controlling COVID-19 in a Hong Kong Chinese Adult General Population. *PLoS ONE*, 17(5), e0268336. Online: <https://doi.org/10.1371/journal.pone.0268336>
- ZHANG, Zhimin et al. (2022): Artificial Intelligence in Cyber Security: Research Advances, Challenges, and Opportunities. *Artificial Intelligence Review*, 55(2), 1029–1053. Online: <https://doi.org/10.1007/s10462-021-09976-0>

A KÖTET SZERZŐI

Bányász Péter (PhD) az ELTE Állam- és Jogtudományi Karán szerzett politológus diplomát, majd a Nemzeti Közszeroláati Egyetem Katonai Műszaki Doktori Iskolájában szerezte meg tudományos fokozatát. Kutatási területei közé tartozik a kiberbiztonság humán aspektusa, a lélektani műveletek hálózatelméleti megközelítése, illetve a generatív mesterséges intelligencia kiberbiztonsági kockázatai. Jelenleg a Nemzeti Közszeroláati Egyetem Államtudományi és Nemzetközi Tanulmányok Karának egyetemi docense, illetve a Kiberbiztonsági Kutatóintézetének kutatója. Több tudományos társaság aktív tagja.

Dub Máté (doktorandusz) a Nemzeti Közszeroláati Egyetem Hadtudományi Doktori Iskolájának hallgatója, kutatási területe a védelmi informatika és kommunikáció elmélete. Kiberbiztonsággal kapcsolatos kutatásait 2018-ban kezdte, amelynek fő területei az adat- és információbiztonság, a biztonságtudatosság, a privátszférát erősítő technológiák, valamint a dezinformációval, a *fake news*-zal szembeni fellépési lehetőségek. Több empirikus vizsgálatot folytatott a *social engineering*/*phishing* támadások és az álhírek terjedésének felismerésének és megelőzésének céljából. A doktori iskolában folytatott kutatói tevékenysége közé tartozik továbbá az új technológiák (például mesterséges intelligencia) jelentette veszélyek és kihívások megjelenése a közösségi média felületeken.

Hankó Viktória (doktoranda) a Nemzeti Közszeroláati Egyetem Katonai Műszaki Doktori Iskolájának hallgatója, 2019 óta foglalkozik a kiberbiztonság különböző aspektusaival. Jártas az adatvédelmi kérdéskörben, az ipari rendszereket érintő kiberbiztonsági kérdésekben, valamint a biztonságtudatosság témakörében is. Doktori munkája során az IoT-rendszerek kritikus infrastruktúrákban levő szerepét vizsgálja, illetve a mesterséges intelligencia megjelenését, kiberbiztonsági kihívásait.

Kovács Gábor (PhD) alkalmazott nyelvész, kommunikáció- és emlékezetkutató. Az NKE Nemeskürty István Tanárképző Karának tudományos dékánhelyettese, a Digitális Média és Kommunikáció Tanszék docense. A Magyar Tudományos Akadémia Bolyai János Kutatási Ösztöndíjával támogatott MAMUT – *Magyar Munkamemória-teszt* fejlesztője. Korábban oktatott a Budapesti Corvinus Egyetemen, a Pannon Egyetemen és az Eötvös Loránd Tudományegyetemen. Fő kutatási területei: az idegennyelv-elsajátítás kognitív háttere, a nyelvi emlékezet, médiareprezentációk, médiahatások és az össztársadalmi jelentőségű viselkedésformák (környezettudatos életmód, online védekezés) motivációs háttere.

Laska Pál (doktorandusz) a Nemzeti Közszołgálati Egyetem Katonai Műszaki Doktori Iskolájának hallgatója, a Gábor Dénes Egyetem tudományos munkatársaként az IT Információbiztonsági Kutatócsoportjában dolgozik. Fő kutatási területe az információbiztonsági tudatosság fejlesztése, a lakossági információbiztonság felmérése és a tudatosítás módszereinek megújítása. Doktori dolgozata az információbiztonsági tudatosságfejlesztés új útjait keresi.

Mezriczky Marcell (doktorandusz) *deepfake*-kutató (BCE) és médiatervező (Wavemaker Hungary). Doktori kutatási témája a technológia társadalomra gyakorolt pozitív-negatív hatásainak vizsgálata. Célja, hogy mélyebb megértést szerezve a területen, hozzájáruljon az „álarcos” videók és képek felismeréséhez, kezeléséhez és kivédéséhez. Emellett médiatervezői háttérének köszönhetően kreatív médiamegoldásokat kíván fejleszteni, amelyek segítségével felismerhetővé és ellenőrizhetővé lehet tenni a *deepfake* tartalmakat.

Szabó Kincső (doktoranda) a Nemzeti Közszołgálati Egyetem NITK Digitális Média és Kommunikáció Tanszék tanársegédje, a Budapesti Corvinus Egyetem kommunikációtudományi doktori hallgatója, a Magyar Diplomáciai Akadémia képzéseinek kurzusvezetője. Kutatási területei közé tartoznak a meggyőzési és előadás-technikák, a toborzói kommunikáció, illetve a nemi sztereotípiák nyelvi vonatkozásai.

Veszelszki Ágnes (PhD) nyelvész, közgazdász, kommunikációkutató. A Nemzeti Közszołgálati Egyetem Nemeskürty István Tanárképző Karának dékánja, az NKE Digitális Média és Kommunikáció Tanszékének tanszékvezető egyetemi docense. A Nemzeti Média- és Hírközlési Hatóság Médiatudományi Intézetében kutatásvezető-helyettes. *A Filológia.hu* MTA-folyóirat főszerkesztője. Korábban a Budapesti Corvinus Egyetem Kommunikáció- és Médiatudomány Tanszékén volt tanszékvezető egyetemi docens, illetve oktatott az Eötvös Loránd Tudományegyetemen, a Budapesti Metropolitan Egyetemen, a Mathias Corvinus Collegiumban és a Szent Ignác Jezsuita Szakkollégiumban is. Főbb kutatási területei: az audiovizuális manipuláció formái, a *deepfake*; a mesterséges intelligencia kommunikációs hatásai; tudománykommunikáció. Pro Scientia témavezetői díjas (2013); Deme László-díjas (2018); a Sánta Zoltán-emlékdíj kitüntetettje (2019). 2020-ban a Corvinus Egyetemen az év oktatójává választották, 2025-ben Az év tudományos diákköri konzulense lett a Nemzeti Közszołgálati Egyetemen.

FÜGGELÉKEK

RÖVIDÍTÉSEK JEGYZÉKE

Rövidítés	Magyar nyelvű kifejezés	Idegen nyelvű kifejezés
AI/MI (a kérdőívben használt rövidítés a 8. és 9. fejezetekben)	mesterséges intelligencia	artificial intelligence
CTF	zászlószerzés	capture the flag
DL	mélytanulás	deep learning
GAN	generatív hálózatok	General Adversarial Networks
MI-6	brit Katonai Hírszerzés 6-os Szekciója	Military Intelligence, Section 6
MFA	többfaktoros hitelesítés	multi-factor authentication
ML	gépi tanulás	machine learning
N	megfigyelés elemszáma	Numerus (latin)
NKE	Nemzeti Közszerzői Egyetem	Ludovika University of Public Service
OECD	Gazdasági Együttműködési és Fejlesztési Szervezet	Organisation for Economic Co-operation and Development
p	p-érték (szignifikancia)	p-value
PMT	védelemmotiváció elmélete	protection motivation theory
VPN	virtuális magánhálózat	virtual private network

A MESTERSÉGES INTELLIGENCIÁVAL
GENERÁLT KÉPI INGEREK ELŐÁLLÍTÁSÁHOZ
HASZNÁLT PROMPTOK

A vizsgálatunkban használt MI-generált képekkel kapcsolatos technikai információk az alábbiak:

AI: Playground

Model: Stable Diffusion

Filter: Realism Engine

Preset: Realistic Portraits (SDXL)

Image Dimension: w:720, h: 1024

Prompt Guidance: 7

Quality& Details: 30

Refinement: 0

Sampler: Euler a

A promptolás során, hogy hatékonyabb szempontokat határozhassunk meg, felhasználtuk a ChatGPT Midjourney – MJ Prompt Generator (V6) nevet viselő alkalmazását, amely különböző finomhangolt promptokat generált a megadott jellemzők alapján. Mindegyik kategória kapcsán két promptot használtunk, és azok segítségével összesen 4-4 képet generáltunk a valódi képekhez hasonló környezetben.

A promptok az alábbiak voltak az egyes kategóriák esetében.

20-as éveiben járó férfi

1–4. kép promptja:

A highly detailed photorealistic image of a young Caucasian male, aged 20, taking a selfie in his room during the evening. He sports a playful expression and a victory hand gesture. He wears a casual hoodie, often worn to rock concerts. His hairstyle is medium-length, diavant style, black. The room features typical youth decor: a bed with colorful bedding, posters of bands,

and scattered personal items. Professional studio lighting emulates a cozy, warm effect. Technical specs: Sony ILCE-7M4, 25.7mm focal length, F/2.8

5–8. kép promptja:

A photorealistic, real-life image of Caucasian man in his 30s. He has his back to the sunset from a hillside. The image is edited in black and white. Technical specifications of the image: dimensions: 4000x6000, aspect ratio: 2:3, camera: Nikon d5600, focal: 140.0 mm, aperture: F/5.6, ISO 200, shutter speed: 1/800. The image is intended to be as photorealistic as possible.

40-es éveiben járó férfi

1–4. kép promptja:

A highly detailed, photorealistic portrait of a 45-year-old man with dark brown hair greying at the temples. He stands with arms folded, upper torso visible, smiling broadly. He's dressed elegantly like an upper-class father. The background shows a living room with soft, natural lighting creating a sunny ambiance, blurred for depth. Professional studio lighting adds a warm, cozy effect. Camera: Nikon D850, 50mm focal, F/3.2, ISO 100, 1/800s shutter speed, aspect ratio: 3:2 —v 6.0

5–8. kép promptja:

Create a photorealistic image of a middle-aged man, 45 years old, with greying dark brown hair. No glasses, no facial hair, his broad smile exudes warmth. He's dressed in sophisticated attire typical of an upper-class father, standing with folded arms. The living room background is softly blurred, bathed in natural sunlight. Camera specs: Nikon D850, 50mm, F/3.2, ISO 100, 1/800s, aspect ratio: 3:2 —v 6.0

60-as éveiben járó férfi

1–4. kép promptja:

A photorealistic depiction of a distinguished 60-year-old Caucasian man in a navy blue suit, sitting in an office environment. His hair is cropped short, he sports a grey beard, and he looks at the camera with a friendly demeanor. The office features a large window with a view of the city centre in the background, the scene illuminated by natural sunlight, creating a soft focus on the distant buildings. Camera settings: Nikon D5600, 140mm, F/5.6, ISO 200, 1/800s, aspect ratio: 2:3 — v 6.0

5–8. kép promptja:

Create a photorealistic image of a Caucasian man in his 60s. The man is wearing a short grey beard, his hair is cropped short, business-like. The man is sitting in an office environment, wearing a classic navy blue suit with matching tie. The man is sitting in front of a laptop computer, giving the impression of a businessman, apparently in a senior position. The man is looking at the camera with a friendly, slightly touching look. In the background is the large window of the office, through which sunlight filters in, with the city centre in the background. The background is slightly blurred. Camera specs: 4000x6000, aspect ratio: 2:3, camera: Nikon d5600, focal: 140.0 mm, aperture: F/5.6, ISO 200, shutter speed: 1/800. The image is intended to be as photorealistic as possible.

20-as éveiben járó nő

1–4. kép promptja:

Create a photorealistic image of a 23-year-old young Caucasian woman. She has long, brunette hair and smiles at the camera. The woman is wearing an autumn fabric jacket, apparently a university student. A backpack is visible on her back. She is holding a book in her hand. The woman is standing in a park, it is a sunny autumn day. The background is blurred faintly. The

composition focuses mainly on her, using natural light to create a welcoming and warm atmosphere. Technical specifications of the image: dimensions: 8256x5504, aspect ratio: 3:2, camera: Nikon D850 camera, focal: 50.0 mm, aperture: F/3.2, ISO 100, shutter speed: 1/800. The image is intended to be as photorealistic as possible.

5–8. kép promptja:

Create a highly detailed, photorealistic portrait of a young 23-year-old Caucasian female student with long brunette hair, smiling at the camera. She's dressed in a seasonal autumn jacket and has a backpack, suggesting a university setting. She holds a book while standing in a sunlit park with autumn leaves. The background is a gentle blur. Nikon D850, 50mm, F/3.2, ISO 100, 1/800s, aspect ratio: 3:2 —v 6.0

40-es éveiben járó nő

1–4. kép promptja:

Create a photorealistic image of a 43-year-old woman with long blonde curly hair. The woman is sitting in a bright, modern living room. She is wearing designer glasses. Her appearance is expressive and elegant, her look is lawyerly. She has a notebook and papers in front of her. The living room has a bookshelf and houseplants. The background shows a living room with soft, natural lighting creating a sunny ambiance, blurred for depth. Camera specs: Nikon D850, 50mm, F/3.2, ISO 100, 1/800s, aspect ratio: 3:2 —v 6.0

5–8. kép promptja:

A highly realistic image of a 43-year-old woman with blonde curly hair, seated in a stylish, modern living room. She wears designer glasses and has an elegant, lawyerly demeanor. In front of her there are a notebook and papers, with a background of bookshelves and houseplants in soft, sunny lighting. The scene is softly blurred for depth. Nikon D850, 50mm, F/3.2, ISO 100, 1/800s, 3:2 aspect ratio —v 6.0

60-as éveiben járó nő

1–4. kép promptja:

Create a photorealistic image of a Caucasian grandmother in her 60s. The woman has long grey hair, tied up in a bun. She is wearing a comfortable dress with a flour covered apron. She is sitting at the table in the kitchen, with freshly baked bread and kitchen equipment in front of her. The kitchen has an old-fashioned, country décor. The woman is looking at the camera with a friendly, slightly touching look. In the background is a window of the kitchen, through which sunlight filters in, with the city centre in the background. Camera specs: 4000x6000, aspect ratio: 2:3, camera: ILCE-7M3, focal: 50.0 mm, aperture: F/4.0, ISO 2000, shutter speed: 0.004s. The image is intended to be as photorealistic as possible.

5–8. kép promptja

Photorealistic portrait of a grandmother in her 60s, Caucasian, with grey hair up in a bun. She's in a kitchen dressed in a comfortable dress and a flour-covered apron, sitting at a table with fresh bread and kitchen utensils. The room is decorated in a country style, with a kitchen window showing sunlight and the cityscape. Friendly expression. Camera: ILCE-7M3, 50mm, F/4.0, ISO 2000, 0.004s, aspect ratio: 2:3 —v 6.0

A VIZSGÁLATBAN ALKALMAZOTT
KÉPEK EREDETI MÉRETBEN

E függelék a 8.2. táblázatban miniatűrként szereplő képi ingereket tartalmazza eredeti méretben, így megfigyelhetők azok a részletek, amelyek alapján ki-töltőink ítéletet alkottak az egyes képek valóságáról. Az online kérdőíven minden egyes kép szélessége egységesen 600 képpont volt. A képeket itt azok bemutatási sorrendjében közöljük.



F.1. ábra: 20-as éveiben járó férfi

Forrás: <https://www.pexels.com/hu-hu/foto/szoba-portre-allo-kep-fuggoleges-kep-17985576/>



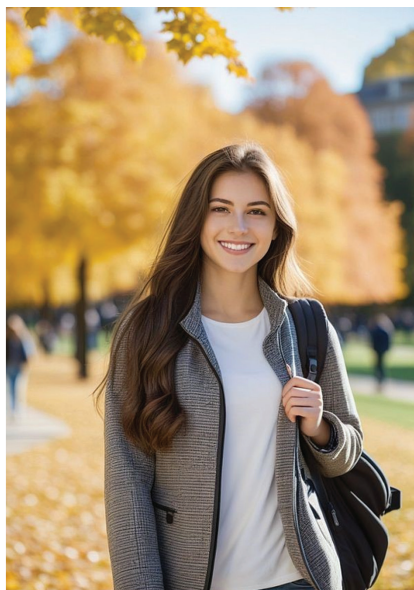
F.2. ábra: 20-as éveiben járó férfi

Forrás: saját szerkesztés Playground AI segítségével



F.3. ábra: 20-as éveiben járó nő

Forrás: <https://www.pexels.com/hu-hu/foto/szemely-no-kave-csesze-712513/>



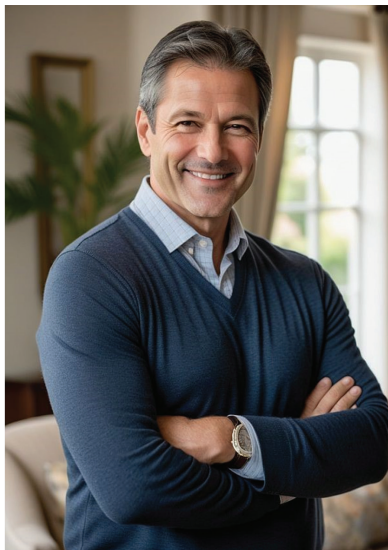
F.4. ábra: 20-as éveiben járó nő

Forrás: saját szerkesztés Playground AI segítségével



F.5. ábra: 40-es éveiben járó férfi

Forrás: https://www.freepik.com/free-photo/smiley-father-posing-with-arms-crossed_6608819.htm#fromView=search&page=1&position=1&uuid=42094065-ddec-476a-a014-36946bb6a67e



F.6. ábra: 40-es éveiben járó férfi

Forrás: saját szerkesztés Playground AI segítségével



F.7. ábra: 40-es éveiben járó nő

Forrás: https://www.freepik.com/free-photo/portrait-expressive-woman_11201718.htm#fromView=search&page=3&position=26&uuid=cca190fb-f9c3-44ef-91db-4535a726a5ed)



F.8. ábra: 40-es éveiben járó nő

Forrás: saját szerkesztés Playground AI segítségével



F.9. ábra: 60-as éveiben járó férfi

Forrás: <https://www.pexels.com/hu-hu/foto/divat-ferfi-szemely-asztal-834863/>



F.10. ábra: 60-es éveiben járó férfi

Forrás: saját szerkesztés Playground AI segítségével



F.11. ábra: 60-as éveiben járó nő

Forrás: <https://www.pexels.com/hu-hu/foto/elelmiszer-zoldsegek-no-ules-4057756/>



F.12. ábra: 60-es éveiben járó nő

Forrás: saját szerkesztés Playground AI segítségével

Hogyan védekezhetünk a digitális tér rejtett fenyegetéseivel szemben?

A közösségi média nemcsak összeköt bennünket, hanem láthatatlan kockázatokkal is szembesít: álhírek, manipulált tartalmak és egyre kifinomultabb mesterséges intelligencia által generált hamis képek próbálják befolyásolni érzékelésünket és döntéseinket.

Ez a tanulmánykötet friss kutatásokon alapulva mutatja be, hogyan érzékelik és kezelik a felhasználók a közösségi médiában jelen lévő veszélyeket, milyen védelmi stratégiák erősíthetik az önhatékonyt, és miként segítheti a tudatos felelősségvállalás az online biztonságot.

A kötet különlegessége, hogy a védelemmotivációs elmélet keretében vizsgálja a közösségi média biztonsági kihívásait, valamint feltárja a generatív mesterséges intelligencia által készített képek felismerésének határait. A bemutatott elemzések újszerű, gyakorlati tanulságokkal szolgálnak mind a tudomány, mind az oktatás, mind pedig a társadalmi prevenció számára.

E könyv mindenkihez szól, aki érteni és kezelni szeretné a digitális korszak kockázatait – kutatókhoz, oktatókhoz, hallgatókhoz és felelős felhasználókhoz egyaránt.