

**EXAMINATION OF MACHINE LEARNING
BASED EMAIL SECURITY APPLIANCES
ROOT CAUSES OF LOW EFFICIENCY
AGAINST SOCIAL ENGINEERING AT-
TACKS****GÉPI TANULÁS ALAPÚ EMAIL BIZTON-
SÁGI MEGOLDÁSOK ALACSONY HATÉ-
KONYSÁGÁNAK GYÖKÉR OKAINAK VIZS-
GÁLATA SOCIAL ENGINEERING TÁMA-
DÁSOKKAL SZEMBEN**KRASNYÁNSZKI Brúnó¹**Abstract**

I have come across many IT solutions that were not able to react effectively to Social Engineering attacks. However, if we examine the process, a malicious email sent by an attacker could be intercepted in many places, depending on the type. For example, I tried to send a malicious email to a business laptop with average protection and I was surprised to find that no alarm was indicated. In my secondary research, I got to know the scientific journal articles related to the deep learning-based detection of Social Engineering in the literature, and after a precise and deep understanding of the technologies, I started to develop a methodology to solve the problem. As my primary research, I conducted interviews about the topicality of the topic and during the interviews I tried to map the current domestic email security situation. Also, I performed measurements on what seemed to me to be suspiciously high-precision models in order to make sure of the accuracy of the results.

Keywords

Machine Learning, AI, ML, Social Engineering, Email Security, Artificial Intelligence

Absztrakt

Nagyon sok olyan informatikai megoldással találkoztam, amik nem voltak képesek megfelelő hatékonysággal reagálni a Social Engineering támadásokra. Pedig amennyiben megvizsgáljuk a folyamatot egy támadó által küldött kártékony emailt sok helyen meg lehetne fogni típustól függően. Példaképpen kipróbáltam egy átlagos védelemmel rendelkező üzleti laptopra egy kártékony emailt küldeni és meglepődve tapasztaltam, hogy semmilyen riasztás nem jelzett. Szekunder kutatásomban megismertem a szakirodalomban a Social Engineering mély tanulás alapú detektálásával kapcsolatos tudományos folyóirat cikkeket és a technológiák pontos és mély megértése után elkezdtem kidolgozni egy módszertant a probléma megoldására. Primer kutatásomként interjúkat folytattam a téma aktualitásáról és az interjúk során megpróbáltam feltérképezni a jelenlegi hazai email biztonsági helyzetet. Valamint a számomra gyanúsán magas pontosságú modelleken végeztem méréseket, hogy megbizonyosodjak az eredmények pontosságáról.

Kulcsszavak

Gépi tanulás, AI, ML, Social Engineering, Email biztonság, Mesterséges Intelligencia

¹ brunokrasnyanszki@stud.uni-obuda.hu | ORCID: 0000-0002-5672-4919 | Junior Researcher, Obuda University Bánki Donát Faculty of Mechanical and Safety Engineering Artificial Intelligence Workshop | Junior Kutató, Óbudai Egyetem Bánki Donát Gépész és Biztonságtechnikai Mérnöki Kar, Mesterséges Intelligencia Műhely

BEVEZETÉS

Ahogy fejlődik az Információ Technológia ezzel együtt fejlődnek biztonsági megoldások és a támadások is. A technológiai kényelmessé teszi az életünket és most már a „technológiai bennszülött” kifejezés is kialakult a szingularitás közeledésével. Ezt kihasználják a támadók is, hiszen ahogy fejlődik az informatikai biztonság a támadóknak új módszereket kellett találniuk és a Social Engineering amik alatt olyan támadásokat értünk amik vagy teljesen pszichológiai vagy pszichológiai és informatikai támadási vektort is tartalmaznak. Ezek a támadások természetükből adódóan a technika számára nehezen detektálható és megfelelő biztonságtudatosság nélkül a célzott támadások közel 100%-os hatékonysággal bírnak. A nem célzott támadásokat nagyobb hatékonysággal lehet detektálni egyszerűségükből fakadóan. Ebből viszont egyre kevesebb lesz meglátásom szerint a későbbiekben mivel az olyan adatszivárgási botrányok, mint például a Cambridge Analytica botrány volt, nagy számú felhasználói adat kerülhetett a dark webre. Ilyen adatok alapján pedig nagyon könnyen lehet célzott támadásokat formálni akár automatikusan is!

Motiváció

Korábbi kutatásaim során (korábbi kutatásom eredményei olvasható a Biztonságtudományi Szemlében jelent meg 2022.09.28-án a 4.évfolyam harmadik számában’ Hogyan változtatta meg a XXI-ik századi kibervédelmet a Social Engineering?’ címmel) is foglalkoztam a Social Engineering jelentette veszélyekkel és az ilyen jellegű támadásokra való felkészüléssel és védekezéssel. Eddigi tapasztalataim alapján kevés cég van azzal tisztában, hogy mekkora veszélyt tudnak jelenteni a pszichológiai támadási vektor is tartalmazó támadások. A támadások mivoltijából fakadóan nehéz ellenük védekezni konvencionális módszerekkel. A műszaki megoldások pedig csak egy részét jelentik az ilyen jellegű védelmi folyamatoknak. Ugyanakkor a helyzetet az is nehezíti, hogy a jelenleg alkalmazott műszaki védelmi megoldások alacsony hatékonysággal működnek az eddig biztonsági vezetőkkel/szakértőkkel készített interjúim alapján. Ez biztatott arra, hogy én is elkezdjek kutatnia témában, mivel nagyon bosszantott, hogy akkori ismereteim szerint nem lehet műszaki eszközökkel hatékonyan védekezni ilyen jellegű támadások ellen.

SOCIAL ENGINEERING TERMINOLÓGIÁK

A Social Engineering kifejezés definíciója kapcsán egy jól ismert szakértő Christopher Hadnagy „The Art of Human Hacking” című könyvében leírt definícióját fogadom el dolgoza-tomban, amely a következő: „A Social Engineering a befolyásolás és rábeszélés eszközével megteveszti az embereket, manipulálja vagy meggyőzi őket, hogy a Social Engineer tényleg az, akinek mondja magát. Ennek eredményeként a Social Engineer – technológia használatával vagy anélkül – képes az embereket információszerzés érdekében kihasználni.” [3] Mivel jelenlegi ismereteim szerint nincs általánosan elfogadott fordítása a magyar szaknyelvben (legtöbbször pszichológiai manipulációt tartalmazó támadásként szokták használni), ezért az angol Social Engineering (röviden SE) kifejezést fogom a továbbiakban használni.

A Social Engineering támadások tárháza nagyon széles, emiatt ebbe a TDK dolgozatomba nem kívánok belemenni az összes támadási típus ismertetésébe. A támadásokat a

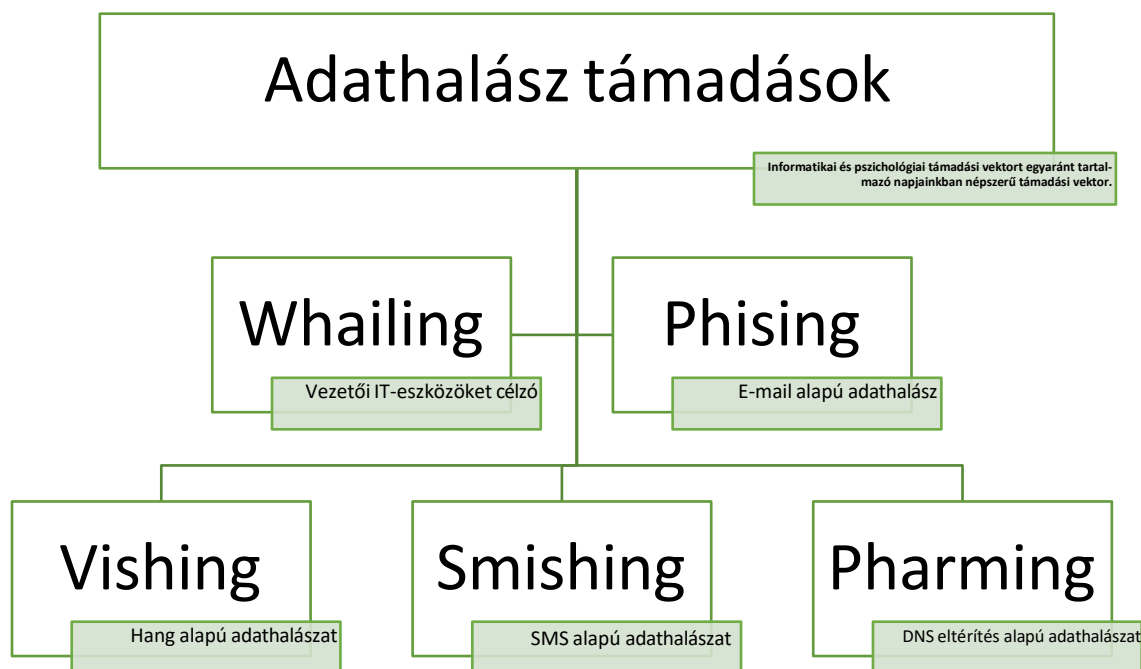
hazai és nemzetközi szakirodalom széles körben részletezi. Nemzetközi szakirodalom kapcsán Christopher Hadnagy fent említett könyve g című könyve, hazai szakirodalom kapcsán pedig Tóth Tamás folyóirat cikke [5] dolgozza fel a Social Engineering támadásokat, valamint én is foglalkoztam vele korábbi kutatásom során [6].

Amivel első körben foglalkozni szeretnék TDK dolgozatomban azok az adathalász támadások. Az adathalász támadások során a támadó valamilyen (jellemzően infokommunikációs csatornákon, de volt példa papír alapú adathalászátra is [7]) módon meggyőzi az áldozatot, hogy adja meg az adatait valamilyen célból olyan legendával, mely nem kelt gyanút az áldozatban. Példa lehet erre a manapság népszerű Netflix-es vagy GLS-es csalások emailben és SMS-ben egyaránt.



1. ábra: szerző által készített példa a Netflixes Smishingre

Az adathalász támadásokat Tóth Tamás módszere alapján klasszifikáltam [5, pp 95-96] a következő csoportokra:



2. ábra: Adathalász támadások (saját szerkesztés)

Ezen belül is első sorban a Phishing-el, másodsorban pedig a Smishing-el a későbbiekben. Ezeket hasonlóan a Social Engineering-hez angolul fogom használni, mivel jelenleg nincs elterjedt és egységes magyar nyelvű kifejezés, ami egy szóban lefedné.

TÉMA AKTUALITÁSA

A Social Engineering alapú támadások évről évre nagyobb számban jelennek meg és a Terrenova Security szerint a 2022-es évben az azonosított phishing támadások 96%-a [2] emailen keresztül érte a szervezeteket, ezért különösen időszerűnek tartottam ezen témával való foglalkozást.

Továbbá személyesen is volt alkalmam rövid interjú erejéig megkérdezni egy ISO 27001-es auditort, egy információbiztonsági vezetőt, egy információbiztonsági tanácsadót, egy információbiztonsági szakértőt, egy telekommunikációs cég üzletágvezetőjét és egy szintén telekommunikációs cég vezető kutató mérnökét a témáról és mind a hat interjú alatt azt a visszajelzést kaptam, hogy az általuk működtetett email biztonsági megoldások meglehetősen primitív elven működnek és nem képesek feltétlen az ígért számokat hozni. Meg erősítettek benne, hogy érdemes foglalkoznom a témával. Interjúim során kiemelném, hogy a megkérdezettek között vegyesen voltak a versenyszférában és a közigazgatásban egyaránt dolgozók, ezzel növelve a kis mintám reprezentációját.

Továbbá megvizsgáltam az egyik legnépszerűbb megoldást, amit phishing támadások ellen tudunk alkalmazni, ez a Cisco Talos Intelligence Group által működtetett PhishTank (phish-tank.org) ami korábban már ismert, többnyire nagy tömeg számára kiküldött adathalász

oldalakat tud azonosítani, amely tovább igazolja azon hipotézisemet, hogy jelenleg a technológia nem képes kiszűrni a célzott támadásokat, pusztán a nagy tömegeket elérő támadás hullámokat.

A téma reprezentációja az MTMT-ben és IEEE Xplore-ban

A Magyar Tudományos Művek Tárának adatbázisából indítottam keresést kulcsszavakra. Ekkor a 'Social Engineering', 'Email biztonság' és 'mély tanulás' kulcsszavakra kerestem rá együttesen, ami nem hozott találatot. A 'Social Engineering' és 'email' szintén nem hozott találatot. Maga a 'Social Engineering'-re viszont már szép számban érkeztek találatok. Összesen 141, melyek között tudományos publikációk és konferencia előadások voltak nagyszámban. 'Email biztonság' és 'mély tanulás' kulcsszavakra szintén nulla találatot adott az MTMT rendszere. Ebből megállapítottam, hogy Magyarországon ez 2023 elején egy kevésbé vizsgált terület.

Az IEEE Xplore adatbázisban is megvizsgáltam a téma reprezentációját. Mind három kifejezést tartalmazó keresést indítottam a fejlett keresési módban (advanced search: ("All Metadata":Social Engineering) AND ("All Metadata":Email security) AND ("All Metadata":Deep Learning))melynek során 11 találatot kaptam vissza az IEEE Xplore kereséséből, ebből 9 konferencia cikk, 1 folyóirat cikk és egy magazin cikk eloszlásban voltak. A fent említett találatok közül egy 2019-es, egy-egy 2020 és 2021-es, a többi pedig 2022-ben készült. Ezáltal arra a megállapításra jutottam, hogy a tématerület nemzetközi kutatása sem jelentős és csak az elmúlt években kezdtek el vele foglalkozni

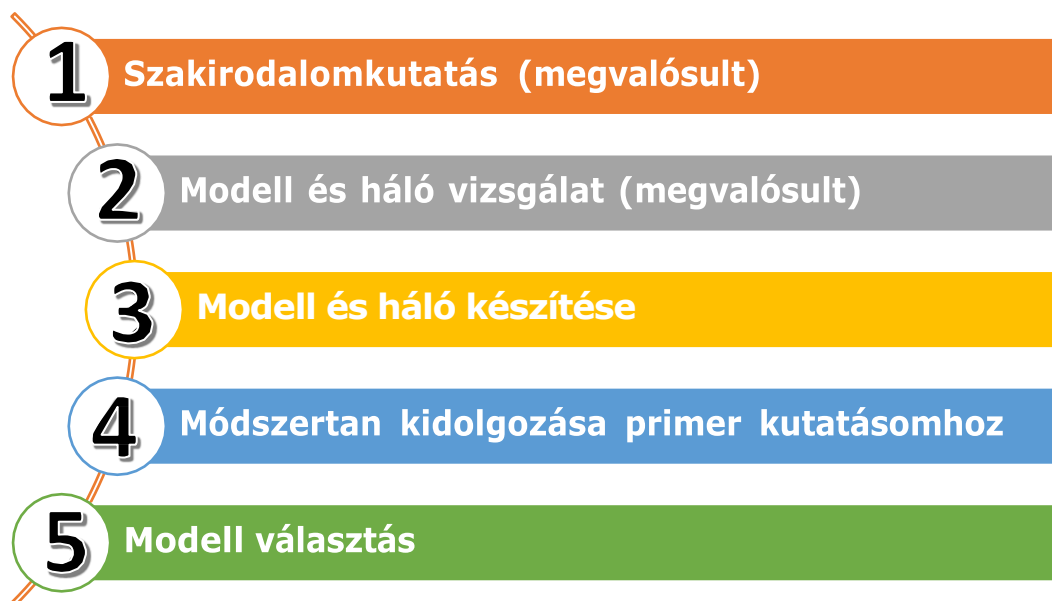
PROBLÉMA MEGFOGALMAZÁSA

Nagyon sok olyan műszaki/informatikai megoldással találkoztam, amik nem voltak képesek megfelelő hatékonysággal vagy egyáltalán detektálni és reagálni a Social Engineering támadásokra. Pedig, ha belegondolunk egy támadó által küldött kártékony emailt sok helyen meg lehetne fogni típustól függően. Példaképpen ki is próbáltam egy átlagos védelemmel rendelkező üzleti laptopra (reaktív vírusirtó, tűzfal, Email Biztonsági Megoldás a levelező szerver előtt) egy kártékony emailt küldeni és meglepődve tapasztaltam, hogy semmilyen riasztás nem jelzett. Továbbá nagy problémának érzem azt, hogy egyre nagyobb arányban alkalmaznak Social Engineering támadásokat a támadók, mivel a konvencionális támadásokra nagyrészt sikertelen műszaki védelmet fejlesztenünk, de sajnos ilyen ütemben nem fejlesztettük a biztonsági tényező leggyengébb láncszemét a humán tényezőt. Ugyanakkor előző kutatásaim és jelenlegi interjúim rávilágítottak arra, hogy nagyon nehéz felkészíteni a munkavállalókat az ellenük professzionális szinten kidolgozott támadásokra, mikor a biztonságtudatosító képzés sokszor egy kimondottan hatékony tűz és munkavédelmi előadás hatékonyságára hajaz. Ezeket persze lehet javítani különböző módszerekkel pl.: gamifikációval [8]. A módszerek közül külön kiemelném Oroszi Eszter Diána biztonságtudatosítási szabadulószoftver megoldását, amivel sokkal hatékonyabban lehet ezen képzéseket elvégezni. Ugyanakkor mivel még mindig az embereket érintő tényezőről van szó, amit nehéz és költséges kezelni, ezért általánosan elmondható, hogy a szervezetek vezetői jobban örülnének olyan műszaki megoldásnak, amit egyszer megvesznek és utána egyáltalán nem, vagy csak kevesebb képzésre kell a munkavállalóit elküldenie.

HIPOTÉZISEK



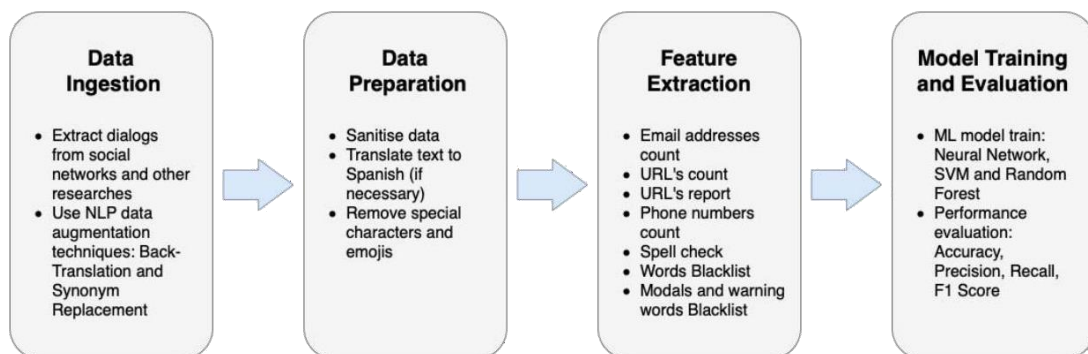
3. ábra: Hipotéziseim (saját szerkesztés)



4. ábra: Szekunder kutatás tervezett lépései (saját szerkesztés)

SZAKIRODALMI ELEMZÉS, KORÁBBAN HASZNÁLT MÓDSZEREK

Szekunder kutatásom alatt a két számomra legígéretesebb modellt fogom ismertetni a fejezetben. Juan C. Lopez és Jorge E. Camargo [14] modellje NLP-t, gépi tanulást, vektoros gép segítséget és véletlenszerű erdő algoritmust alkalmaz. Ezzel állításuk szerint 84%-os pontosságot tudtak elérni. A rendszerük grafikusán a következő képpen néz ki:



5. ábra [14, p 178, fig1]

Ennél a modellnél nagyon előre mutatónak érzem, hogy NLP-vel már bizonyos Social Engineering támadási formákat már képes felismerni, mint pl.: ráhatás, túlterhelés és a tekintély bizonyos fajtáit. [14, p 179] Ezen felül a konvencionális eszköztárból is használtak URL vizsgálatot, szó fekete listát és telefonszám vizsgálatot.

Bronjon Gogoi és Tasiruddin Ahmed [15, p 2] modellje transfer learning-et alkalmaz annak érdekében, hogy kevesebb adatból is képes legyen pontosabb döntéseket hozni. A modell alapjául a Google által fejlesztett BERT (Bidirectional Encoder Representation from Transformers) és DISTILBERT modelleket használták, mivel ezekre nagy előre tanított corpusok voltak elérhetőek. Használtak továbbá NLP-t, amit NSP-vel (Next Sentence Prediction – Következő Mondat Jóslása) és NER (Named Entity Recognition - Nevezett Entitás Felismeréssel). A modellük grafikus szemléltetése:

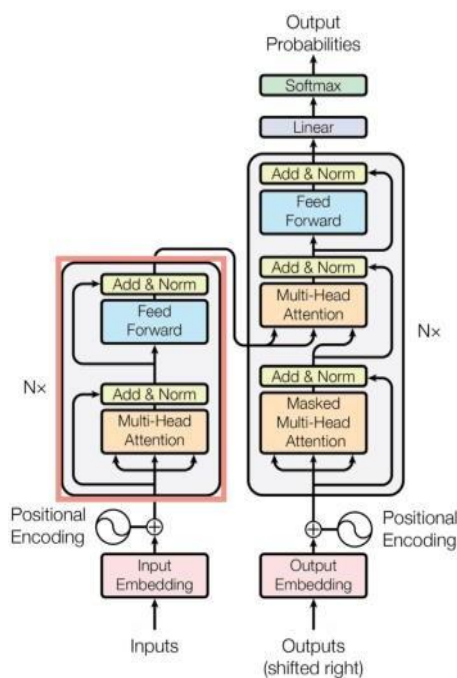


Fig. 3. Transformer Architecture

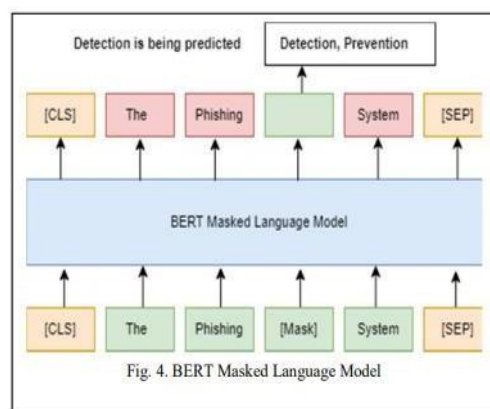


Fig. 4. BERT Masked Language Model

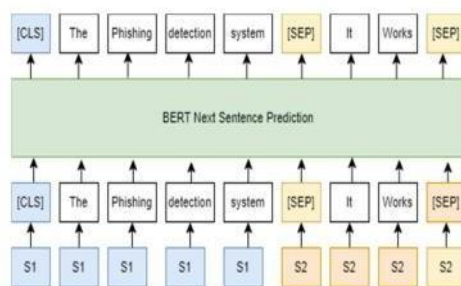


Fig. 5. BERT Next Sentence Prediction Model

6. ábra [15]

PRIMER KUTATÁS

Primer kutatásom jelenleg még kezdeti fázisban tart. Interjúkat készítettem és a szekunder kutatásom során látottakhoz készítettem méréseket, hogy megbizonyosodjak olyan adatokban, amikben nem bízam.

Interjúk az email biztonságról

A készített interjúk első része inkább a témám irányát határozta meg, a második része pedig olyan kihívásokat vetett fel, amik jelenleg még nem megoldottak. Az egyik leghasznosabb interjút egy telekommunikációs-szolgáltató cég üzletágvezetőjével sikerült készítenem, ahol betekintést kaphattam a telefonszolgáltatások biztonsági architektúrájába. Ezen interjú alkalmával rákérdeztem arra is, hogy az 1. ábrámban említett pl.: Netflixes csalások ellen milyen védekezési mechanizmusokat használnak és mi az oka, hogy ennyire elterjedtek az SMS-es adathalás támadások. Azt a meglepő választ kaptam, hogy két védekezési mechanizmust használ a szervezete (ugyanakkor kiemelte, hogy nem csak ők, hanem a többi Magyarországi szolgáltató is). Az első egy olyan empirikus predikció, amely gyanúsán sok kiküldött SMS esetén beriaszt és jelzi a szolgáltatónak, hogy pl.: 500.000 SMS-t küldtek ki ugyanazzal a szöveggel. A másik védelmi mechanizmus SMS-ek szűrésére az a jelentés alapú telefonszám kitiltás. Ez esetben, amikor kellő számú bejelentés érkezik a szolgáltatóhoz és kivizsgálják a gyanús eseteket akkor ideiglenesen letiltatják a telefonszám SMS küldési lehetőségét a hálózatról. Ezen túl viszont mást nem alkalmaznak. Ezzel a szolgáltatások biztonságát érintő megosztott felelősségi modell szerint²² több felelősség helyeznek a felhasználóra mind a lakossági ügyfelek, mind a vállalatok esetén. Amennyiben viszont a vállalatoknál nincs megfelelő hatékonyságú biztonság-tudatosító képzés ezeket a támadásokat nagy eséllyel nem fogják tudni kivédeni.

A többi interjú tartalma és eredménye megegyezett. Interjú alanyaimmal email biztonságról beszélgettem és minden interjún azt az eredményt kaptam vissza, hogy a phishing emaileket sehogyan nem tudják megfelelően szűrni. Egy alkalommal merült fel, hogy megoldás lehet, ha az email szabályzatban letiltjuk a szervezeten kívüli emailek fogadását ezzel tovább erősítve a Zero Trust³ megközelítést, de ez csak kevés helyen vitelezhető ki, mivel a vállalatoknak nem a belső levelezésük termel majd bevételt.

Modellek tesztelése

Sokszor futottam bele primer kutatásom során abba, hogy nagyon magas hatékonyságúnak ígérkező klasszifikáló modelleket találtam a közösség által fejlesztett Kaggle nevű oldalon. [10] Itt sokszor 90+ % hatékonyságú modelleket láttam, viszont nekem ez gyanúsnak tűnt, mivel amikor korábban készítettem egy interjút egy telekommunikációs cég kutató mérnökével, tőle azt az információt kaptam, hogy annak ellenére, hogy piacvezető beszállítók ilyen technológiában viszont 30-40%-nál nagyobb hatékonyságot nem tudtak

² A megosztott felelősségi modell (angolul: Shared Responsibility Modell) a felhő biztonságban elterjedt biztonsági kifejezés, de véleményem szerint más kiberbiztonsági területeken is remekül használható különböző kockázati tényezőinek felírására a sok tényezős biztonsági kihívások elemzésére. [9]

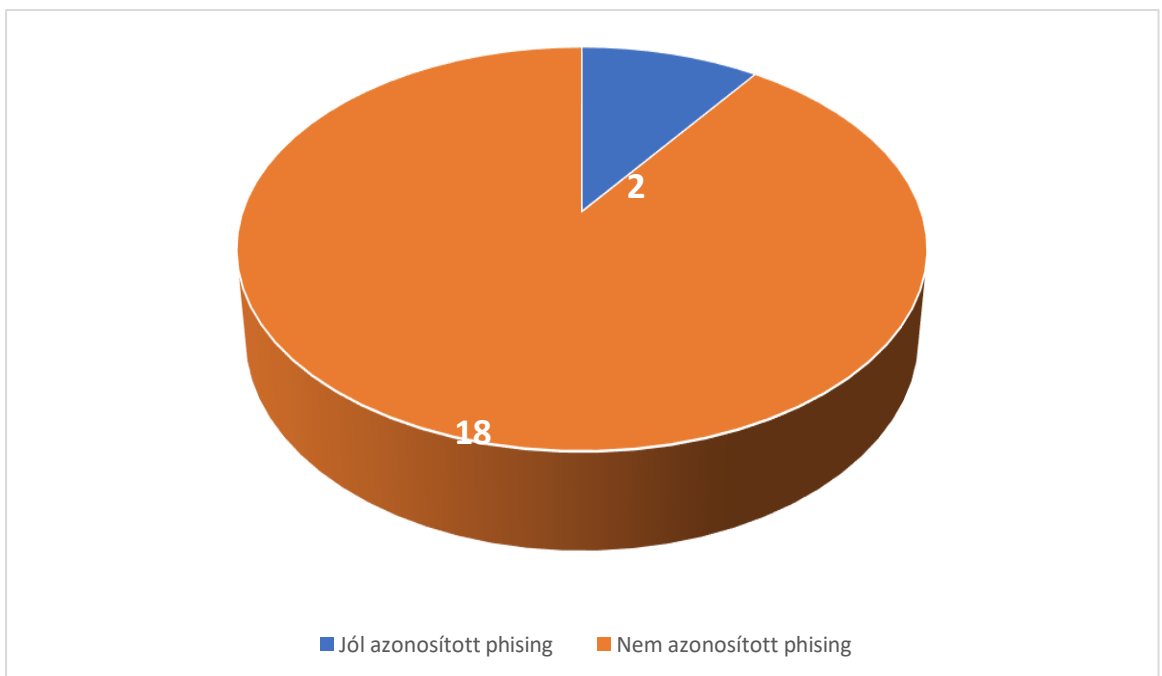
³ Bizalmatlanság elve. Egy olyan kiberbiztonsági jó gyakorlat a jogosultság kezelésnél, amivel csak az kap jogosultságot az adott folyamathoz vagy erőforráshoz, akinek feltétlenül szüksége van rá. Ezzel megelőzhetővé válik, hogy a fertőzés a szervezeten belül elterjedjen, illetve a visszaélések is visszaszoríthatóak vele.

elérni.

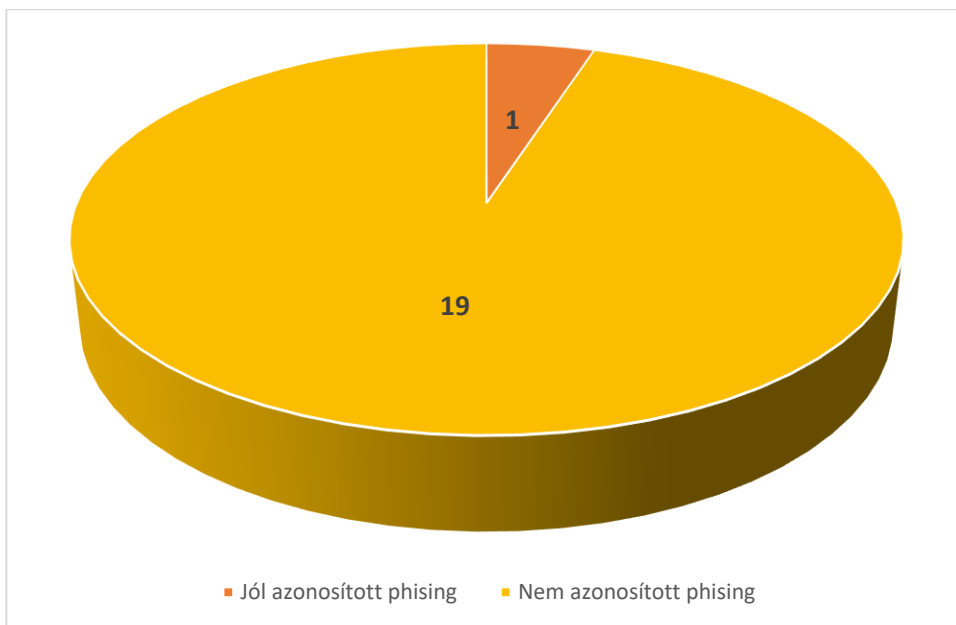
Ebből arra következtettem, hogy ahol nagyon magas hatékonyságot látok, ott a modellt túltaníthatták, ennek következményében az eredeti adat halmazával jó eredményeket tud elérni, viszont valós támadások esetén már sokkal kevésbé.

Ennek igazolására 3 modellnél [11][12][13] megmértem mi történne, ha én írnék phishing emaileket, amiket a papíron jól teljesítő modelleknek kellene klasszifikálnia. Mind három modell esetében 20-20 emaillel végeztem méréseket. Fontos azonban kiemelni, hogy a használt modellek alapvetően binárisan foglalkoztak az email biztonsággal. Vélhetően az egyszerűség és vagy a hatékonyság (futási időre értendő, nem pontosságra) kedvéért, de nem tesznek különbséget a SPAM, phishing vagy egyéb más kártékony emailek között. A

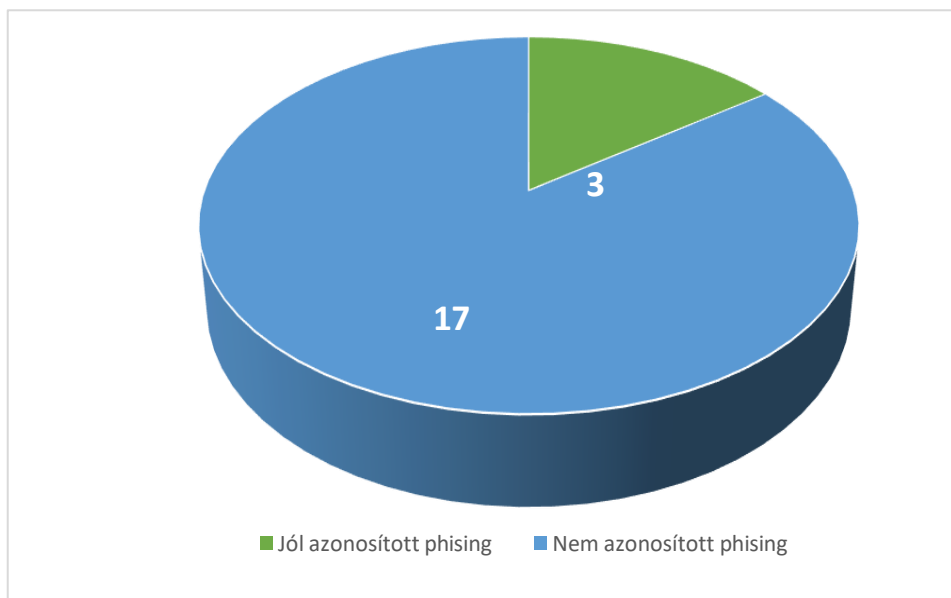
SPAM kategóriába kerül minden kártékony email. Ezt az eredmények értékelésénél fontos figyelembe venni. Az eredményeim grafikus formában itt láthatóak:



7. ábra: Az első modell (saját szerkesztés)



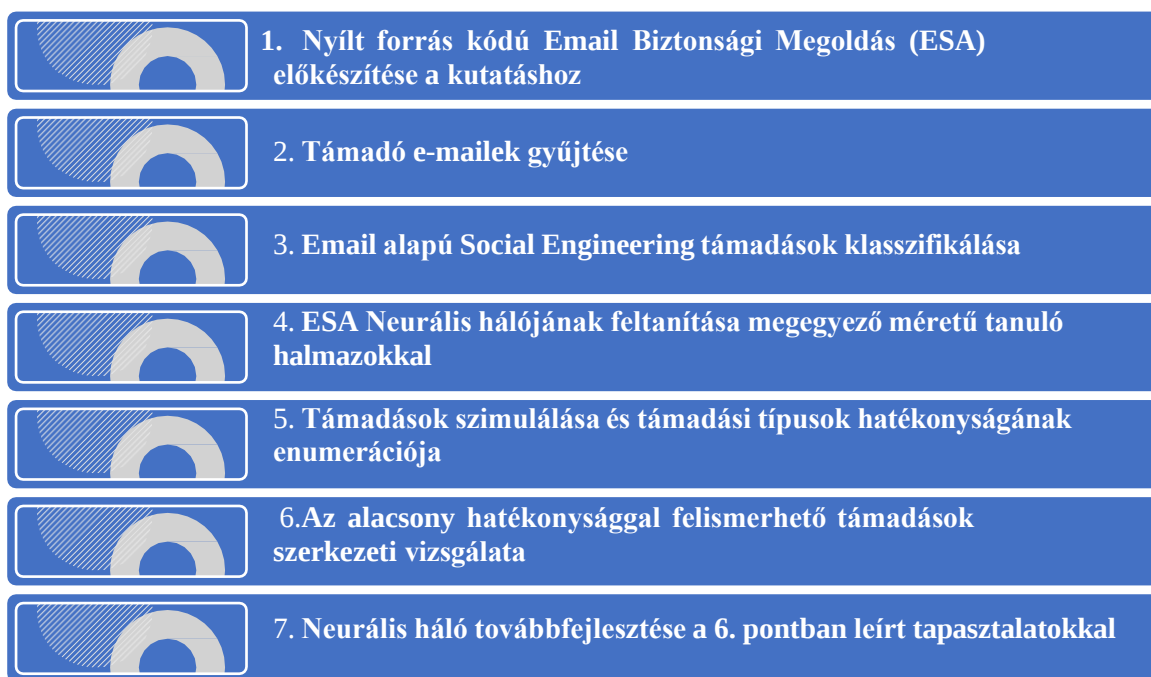
8. ábra: A második modell (saját szerkesztés)



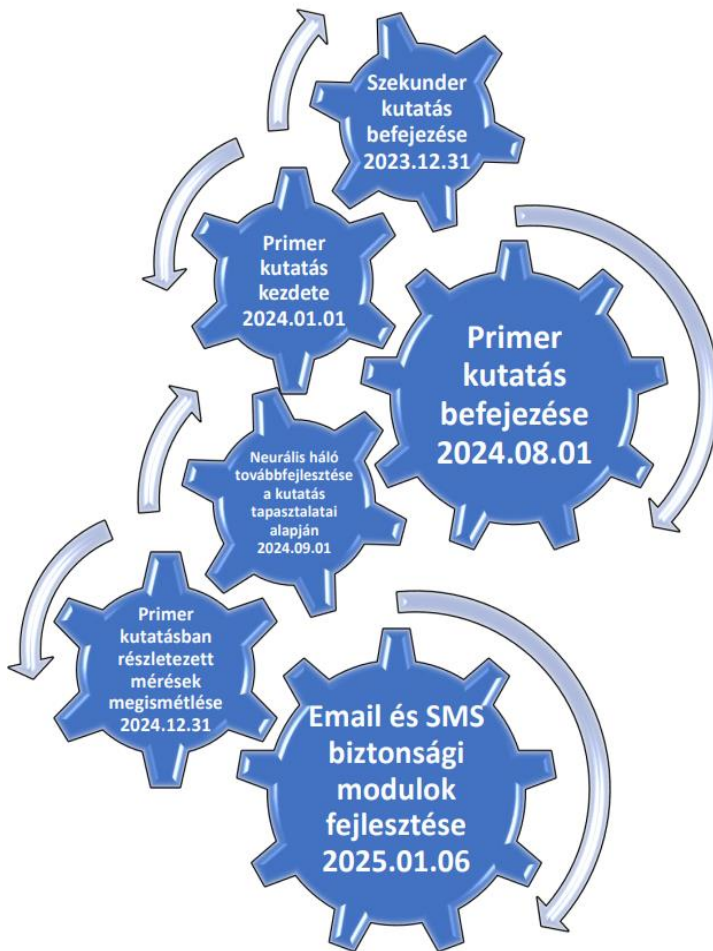
9. ábra: A harmadik modell (saját szerkesztés)

PRIMER KUTATÁSI TERV

Primer kutatásom során szeretném megvizsgálni a legnépszerűbb nyílt forrás kódú Email Biztonsági Megoldások alacsony hatékonyságának okait azzal, hogy különböző Social Engineering támadásokat fogok tesztelni az előzőleg betanított Email Biztonsági Megoldáson igyekezve arra, hogy a tanítóhalmazomban minden típusú támadásból nagyjából ugyanolyan arányban tartalmazzon és ezek után minden típusú támadási típust egyesével tesztelve, ezzel eljutva a mély tanuló algoritmus számára a legkevésbé érzékelhető támadásig, amit ezek után vizsgálatnak vetnék alá, hogy mi volt a gyökér ok, ami miatt nem jelezte a mély tanuló algoritmus gyanúsnak a támadó emailt.



10. ábra: Kutatási terv (saját szerkesztés)



11. ábra: Időterv (saját szerkesztés)

KONKLÚZIÓ

Mivel kutatásom korai szakaszban van ezért jelenleg nem tudtam még mélyreható konklúziókat levonni.

A H1-es hipotézisemet bizonyítottam az interjúk eredményei és a saját méréseim alapján is.

A H3-as hipotézisem a szakirodalomban bizonyítást nyert, ugyanakkor a jövőben szeretném elvégezni a saját méréseimet is.

A H2, H4, H5 és H6-os hipotézisek jelenleg még nem eldönthetőek a rendelkezésre álló információkból.

IRODALOMJEGYZÉK – FORRÁSOK

- [1] ACM CSS Téma klasszifikáció az OTDT követelmények szerint az érintett témából generálva: <https://dl.acm.org/ccs> (generálás ideje: 2023. 05. 01 22:43)
- [2] <https://terranovasecurity.com/cyber-security-statistics/> (hozzáférés: 2023. 05. 02. 14:44)
- [3] Christopher Hadnagy „The Art of Human Hacking” ISBN: 978-1-118-02801-8 WileyPublishing, Inc. (2011) p.23 szerző által fordított
- [4] Kevin D. Mitnick „A legendás hacker - A megtévesztés művésze” 9789632065557
- [5] TÓTH TAMÁS „AZ EGYES SOCIAL ENGINEERING MÓDSZEREK ELHATÁROLÁSA ÉS RENDSZEREZÉSE”, Katonai Nemzetbiztonsági Szolgálat Szakmai Szemle XVIII. évfolyam 1. szám 2020. március HU ISSN 1785-1181
- [6] Krasnyánszki Brúnó „Hogyan változtatta meg a XXI-ik századi kibervédelmet a Social Engineering?”, Évf. 4 szám 3 (2022): Biztonságtudományi Szemle Információbiztonság rovat
- [7] Dub Máté, Szakdolgozat, 'A kiberbiztonság humán aspektusának vizsgálata social engineering technikával', Nemzeti Közszerzői Egyetem Államtudományi és Nemzetközi Tanulmányok Kar Közszerzési és Infotechnológiai Tanszék Közigazgatás-szervező szak Nemzetközi közigazgatási szakirány pp, 35-50
- [8] Oroszi Eszter Diána, Biztonságtudatossági szabadulószoa, mint új programelem az információbiztonsági képzésekben Dunakavics A Dunaújvárosi Egyetem online folyóirata 2020. VIII. évfolyam III. szám pp 51-62
- [9] Shared Responsibility Modellről lásd a Felhő Biztonsági Szövetség megfogalmazását: <https://cloudsecurityalliance.org/blog/terms/shared-responsibility-model/> (utolsó hozzáférés 2023. 05. 03 12:24:38)
- [10] <https://www.kaggle.com/> (utolsó hozzáférés 2023. 05 .03 12:25:21)
- [11] <https://www.kaggle.com/code/mfaisalqureshi/email-spam-detection-98-accuracy> (utolsó hozzáférés 2023. 05 .03 14:13:32)
- [12] <https://www.kaggle.com/code/harshsinha1234/email-spam-classification-nlp>(utolsó hozzáférés 2023. 05 .03 14:14:43)
- [13] <https://www.kaggle.com/code/balaka18/email-spam-classification> (utolsó hozzáférés 2023. 05 .03 14:15:53)
- [14] Juan C. Lopez, Jorge E. Camargo, 'Social Engineering Detection Using Natural Language Processing and Machine Learning' 2022 5th International Conference on Information and Computer Technologies (ICICT), DOI 10.1109/ICICT55905.2022.00038
- [15] Bronjon Gogoi, Tasiruddin Ahmed, ' Phishing and Fraudulent Email Detection through Transfer Learning using pretrained transformer models', 2022 IEEE 19th India Council International Conference (INDICON) | 978-1-6654- 7350-7/22/\$31.00 ©2022 IEEE | DOI: 10.1109/INDICON56171.2022.10040097