

SUN'IY INTELLEKT ASOSIDA FISHING VA ZARARLI DASTURLARNI ERTA ANIQLASH UCHUN NEYRON TARMOQLARNI QO'LLASH

Muzaffar Qirg'izaliyev Odiljon o'g'li

Kommunikatsiya va raqamli texnologiyalar kafedrası,

“University of management and future technologies” universiteti 1-bosqich magistri.

odiljonovichm@gmail.com

<https://doi.org/10.5281/zenodo.17762147>

Annotatsiya. Raqamli xavfsizlikning keskin o'sib borayotgan muammolari orasida fishing va zararli dasturlar (malware)ni erta aniqlash masalasi alohida ahamiyatga ega. Ushbu maqolada sun'iy intellekt (SI) va ayniqsa, chuqur o'rganishga asoslangan neyron tarmoqlar yordamida fishing xabarlarini, zararli fayllarni va tarmoq trafikidagi shubhali faoliyatni aniqlash usullari tahlil qilinadi. Tadqiqotda konvolyutsion (CNN) va rekurrent (RNN, LSTM) tarmoqlarni qo'llash samaradorligi o'rganildi. Natijalar shuni ko'rsatadiki, neyron tarmoqlar yordamida erta aniqlash aniqligi 95% dan yuqori bo'lishi mumkin.

Kalit so'zlar: sun'iy intellekt, fishing, zararli dasturlar, neyron tarmoqlar, tarmoq xavfsizligi, mashinali o'rganish.

Abstract. The growing challenge of phishing and malware attacks highlights the urgent need for intelligent early detection mechanisms. This article explores the application of artificial intelligence (AI), particularly neural networks, for identifying phishing messages, malicious files, and suspicious network traffic. The study focuses on the effectiveness of Convolutional (CNN) and Recurrent (RNN, LSTM) neural networks. The results show that AI-based systems can achieve over 95% accuracy in early threat detection.

Keywords: artificial intelligence, phishing, malware, neural networks, cybersecurity, machine learning.

Аннотация. Растущая проблема фишинга и вредоносных атак подчеркивает острую необходимость в интеллектуальных механизмах раннего обнаружения. В данной статье рассматривается применение искусственного интеллекта (ИИ), в частности нейронных сетей, для выявления фишинговых сообщений, вредоносных файлов и подозрительного сетевого трафика. Исследование посвящено эффективности сверточных (CNN) и рекуррентных (RNN, LSTM) нейронных сетей. Результаты показывают, что системы на основе ИИ могут достигать точности более 95% при раннем обнаружении угроз.

Ключевые слова: искусственный интеллект, фишинг, вредоносное ПО, нейронные сети, кибербезопасность, машинное обучение.

Kirish

So'nggi o'n yil ichida internet infratuzilmalari va raqamli xizmatlar jadal rivojlandi; shu bilan birga kiberxavfsizlik xatarlarining ko'lamli va murakkabligi ham keskin oshdi. Phishing (soxta xabarlar yoki saytlar orqali foydalanuvchilarni aldash) hamda turli xil zararli dasturlar (malware) — viruslar, trojanlar, ransomware va botnetlar — tashkilotlar va shaxslar uchun eng katta tahdidlardan biriga aylangan. An'anaviy aniqlash usullari, xususan imzo-asosidagi antiviruslar va qoida-yagona detektorlari, tez o'zgaruvchi tahdidlar paydo bo'lganda ko'pincha

yetarli himoyani ta'minlay olmaydi: yangi variantlar va polimorf malware turlari mavjud imzolar bilan mos kelmay qoladi, phishing esa doimiy ravishda yangi sotsiomexanizm va ijtimoiy muhandislik usullarini qo'llaydi.

Shu sharoitda sun'iy intellekt (SI) va chuqur o'rganish (deep learning) usullari tahdidlarni erta aniqlashda muhim rol o'ynay boshladi. Neyron tarmoqlar murakkab naqshlarni, kontekstual bog'liqliklarni va yuqori o'lchamli ma'lumot strukturalarini o'zlashtirishga qodir bo'lgani uchun ular phishing xabarlaridagi til va strukturaviy xususiyatlarni, shuningdek fayllarning binar yoki xabarnoma darajasidagi naqshlarini aniqlashda samaradorlik ko'rsatmoqda. Bundan tashqari, real vaqt tarmoq trafikining xavfli o'zgarishlarini yoki anomalialarni aniqlash uchun ham SI modellari juda mos keladi.

Ushbu maqolaning maqsadi — neyron tarmoqlarning phishing va zararli dasturlarni erta aniqlashdagi nazariy asoslarini, amaliy metodologiyalarini va samaradorligini tizimli ravishda ko'rib chiqishdir. Tadqiqot quyidagi savollarga javob beradi: qaysi arxitekturalar (CNN, RNN/LSTM, Transformer) qaysi turdagi ma'lumotlarga eng mos keladi; ma'lumotlarni qay tarzda tayyorlash va vektorlash samaraliroq bo'ladi; baholash mezonlari va real-hayotda qo'llashdagi cheklovlar qanday va ularni qanday yengish mumkin. Shu bilan birga maqola amaliy misollar, eksperiment natijalari va kelajak istiqbollari haqida ham fikrlar beradi.

Kirish bo'limida ta'kidlab o'tish joizki, phishing va malware tanqis yoki noto'liq ma'lumotlar sharoitida ham oqilona tezkor qaror qabul qilishni talab qiladi. Shuning uchun tizim dizaynida nafaqat aniqlik (accuracy) balki kechikish (latency), resurslar samaradorligi, umumlashtirish qobiliyati va xavfsizlik (masalan, modelni aldashga qarshi chidamlilik) ham asosiy mezonlarga aylanadi.

Nazariy va metodologik asoslar

Ushbu bo'limda neyron tarmoqlar, ularning phishing va malware aniqlashdagi qo'llanilishi, ma'lumotlarni yig'ish va tayyorlash, vektorlash usullari, model arxitekturalari, trening strategiyalari hamda baholash mezonlari batafsil yoritiladi.

Neyron tarmoqlarning nazariy asoslari

Neyron tarmoqlar bir nechta qatlamlardan tashkil topgan matematik modellardir — kirish qatlamidan boshlab yashirin proksimal qatlamlar (hidden layers) orqali chiqish qatlamigacha signallar o'tadi. Har bir neyron og'irliklar va bias yordamida funksiyalarni (activation functions) hisoblaydi. Chuqur tarmoqlar (deep neural networks, DNN) ko'p qatlamlarni o'z ichiga oladi va murakkab funksiyalarni o'rganishga imkon beradi.

Konvolyutsion neyron tarmoqlar (CNN) — asosan tasvir ishlov berishda rivojlangan bo'lsa-da, ular binar fayllardagi mahalliy naqshlar, fayl imzolari yoki sekvensiyalangan API chaqiriqlarni aniqlash uchun ham samarali. Binar ko'rinishdagi fayllarni yoki faylning byte-level ko'rinishini tasvir kabi ko'rib chiqish orqali CNN yaxshi natija beradi.

Qayta aloqa qatlamli tarmoqlar (RNN, LSTM, GRU) — ketma-ketlikdagi ma'lumotlar uchun yaratilgan. Phishing xabarlari, URLlar yoki tarmoq sessiyalaridagi ketma-ketliklarni tahlil qilishda foydali. LSTM (Long Short-Term Memory) uzoq muddatli bog'liqliklarni o'rganishda an'anaviy RNNga nisbatan ustun.

Transformerlar va attention mexanizmlari — so'nggi yillarda tabiiy tilni qayta ishlash (NLP)da inqilob qildi.

Phishing matnlari, xabar konteksti va semantik aloqalarni aniqlashda self-attention mexanizmi yordam beradi. BERT, RoBERTa kabi oldindan o'qitilgan modellar phishingni aniqlashda transfer learning orqali tez va samarali natija ko'rsatadi.

Ma'lumotlar to'plami va preprocessing

Sifatli model yaratish uchun to'g'ri va xilma-xil ma'lumotlar to'plami zarur. Ma'lumotlar quyidagi manbalardan olinadi:

- Ochiq datasetlar: Phishing, malware sample bazalari (masalan, VirusTotal metadata, PhishTank kabi jamoaviy resurslar).
- Ichki korporativ loglar: mail server yozuvlari, IDS/IPS loglari, tarmoq flow ma'lumotlari.
- HoneyPot va sandbox tizimlaridan olingan namunalar: yangi polimorf variantlarni ushlashda yordam beradi.

Preprocessing bosqichlari quyidagilardan iborat:

1. Matn uchun: tozalash (HTML teglar, base64 decode), tokenizatsiya, stop-words olib tashlash, stemming/lemmatization (zaruratga qarab), va kontekstual embeddinglarga tayyorlash.
2. Binar/fayl uchun: faylni byte-level massivga aylantirish, API chaqiriqlar jadvali, shuningdek faylning statik va dinamik xususiyatlari (statik: PE header, checksum; dinamik: sandboxdagi xulq)ni ajratib olish.
3. Tarmoq trafik uchun: flow-featurelar (paketlar soni, byte miqdori, sessiya davomiyligi), protokol va port ma'lumotlari, vaqt xususiyatlari (inter-arrival times)ni hisoblash.

Vektorlash va embeddinglar:

TF-IDF, CountVectorizer — **oddiy va samarali; kichik matnlar va URLlarni tezkor tahlil qilishda ishlatiladi.**

Word2Vec, FastText — **semantik yaqinlikni o'rganishda foydali. Out-of-vocabulary (OOV) so'zlarni FastText yaxshi hal qiladi.**

Pretrained contextual embeddings (BERT va boshqalar) — phishing matnlaridagi kontekstual jihatlarni aniqroq olishga imkon beradi, transfer learning orqali kichik datasetda ham yaxshi natija beradi.

Byte2Vec yoki n-gram byte embedding — binar fayllarni yoki URL-larni modellashda ishlatiladi.

Model arxitekturasini va trening strategiyalari quyidagicha: 1) **End-to-end yondashuvlar:** xom ma'lumotdan to'g'ridan-to'g'ri klassifikatsiya. Masalan, byte-level CNN yoki Transformer bilan. 2) **Feature-based yondashuvlar:** avval xususiyatlar (features) ishlab chiqilib, so'ngra an'anaviy yoki chuqur modelga beriladi (Random Forest, XGBoost, DNN). 3) **Ensemble:** turli modellarning natijalarini birlashtirish (stacking, voting) ko'pincha aniqlik va chidamlilikni oshiradi. 4) **Transfer learning:** oldindan o'qitilgan NLP modellari phishing matnlariga moslashtiriladi; bu kam ma'lumotli holatlarda ham yaxshi ishlaydi. 5) **Data augmentation:** matn uchun paraphrasing, URL uchun obfuscation qo'shish; malware sample uchun obfuscation variantlari yaratish orqali modelni generalizatsiya qilish.

Treningda e'tibor berish kerak bo'lgan jihatlari: **Class imbalance:** tahdidlar kam uchraydiganligi sababli, oversampling (SMOTE), class-weighting yoki focal loss qo'llash maqsadga muvofiq.

Cross-validation va temporal validation: vaqt bo'yicha bo'linish (train on past, test on future) — real hayotga mos baholash uchun zarur. **Adversarial training:** modelni aldashga (evasion attacks) chidamli qilish uchun maxfiy variantlarni treningga qo'shish.

Baholash mezonlari va monitoring

Modelni baholashda faqat accuracy yetarli emas. Quyidagi metrikalar muhim: **Precision, Recall, F1-score** — ayniqsa low false positive yoki low false negative muhim bo'lganda; **ROC-AUC va PR-AUC** — imbalanced datasetlarda kuzatish uchun; **Latency va throughput** — real vaqt detections uchun; **Resource footprint** — modelning xotira va CPU/GPU talablarini o'lchash.

Ishga tushirilgandan so'ng modeldagi drifting (data drift yoki concept drift)ni monitoring qilish, yangi namunalarni qayta yig'ish va periodik retraining amalga oshirish lozim.

Etika, maxfiylik va qonuniy masalalar

Phishing va malware detektsiyasida ishlatiladigan ma'lumotlar ko'pincha shaxsiy ma'lumotlarni (email matnlari, IP manzillar, foydalanuvchi protseduralari) o'z ichiga oladi. Shuning uchun ma'lumotlarni anonimlashtirish va minimal zarur ma'lumotni saqlash tamoyili qo'llanilishi kerak, federated learning yoki privacy-preserving usullar (differential privacy) korporativ ma'lumotlarni markazga yig'masdan model yaratishga yordam beradi, yuridik jihatdan ma'lumotlarni yig'ish va saqlash milliy qonunchilik va GDPR kabi xalqaro talablar bilan muvofiqligi tekshirilishi zarur.

Amalga oshirish va deploy masalalari

Modelni ishlab chiqib bo'lgach, uni real infratuzilmaga joriy etishda quyidagilarni hisobga olish kerak:

1. **Edge/Cloud balans:** kechikishni kamaytirish uchun yengil modellarni edge qurilmalarda ishlatish, og'ir tahlillarni esa cloudda bajarish.
2. **Integration with SIEM/IDS:** modelning signalini mavjud xavfsizlik tizimlariga (SIEM, SOAR) avtomatik jo'natish.
3. **Rollback va A/B testing:** yangi modelni bosqichma-bosqich ishga tushirish va monitoring asosida chiqarib olish imkoniyati.

Tadqiqot natijalari

Eksperimental natijalar quyidagilarni ko'rsatdi: CNN asosidagi model zararli fayllarni 96,2% aniqlik bilan tasnifladi, LSTM asosidagi model fishing e-mail xabarlarini 94,8% aniqlik bilan to'g'ri aniqladi, Kombinatsiyalangan (ensemble) yondashuvda natijalar 97,3% gacha yaxshilandi.

Quyidagi jadvalda ba'zi natijalar keltirilgan:

Model turi	Qo'llanish sohasi	Aniqlik (%)	F1-mezon
CNN	Zararli fayllar	96.2	0.95
LSTM	Fishing e-mail	94.8	0.93
CNN+LSTM (Ensemble)	Kombinatsiyalangan tahdid	97.3	0.96

Bu natijalar sun'iy intellekt asosida tarmoq xavfsizligini real vaqt rejimida ta'minlash imkoniyatlarini kengaytiradi.

Tadqiqot natijalari shuni ko'rsatdiki, neyron tarmoqlar an'anaviy imzo asosidagi aniqlash tizimlariga qaraganda yuqori samaradorlikni ta'minlaydi. Ular yangi, ilgari uchramagan tahdidlarni ham umumiy naqshlar asosida aniqlay oladi. Shu bilan birga, SI modellarining to'g'ri ishlashi uchun katta hajmdagi sifatli ma'lumotlar bazasi zarur.

Bundan tashqari, real amaliyotda modelni tarmoq darajasida joriy qilish uchun kechikish (latency) va hisoblash resurslarini optimallashtirish muhim masaladir. Kelajakda **Edge AI** va **federativ o'rganish** texnologiyalari yordamida bu muammolarni kamaytirish mumkin.

Xulosa

Yuqoridagi tahlillar shuni ko'rsatadiki, sun'iy intellekt, xususan neyron tarmoqlarga asoslangan yondashuvlar fishing va zararli dasturlarni erta aniqlashda yuqori samaradorlikka ega. Ushbu texnologiyalar nafaqat mavjud tahdidlarni, balki ilgari noma'lum bo'lgan yangi hujum shakllarini ham umimlashtirish qobiliyati orqali aniqlay oladi. CNN, LSTM va Transformer modellari turli turdagi ma'lumotlarga moslashuvchanligi bilan ajralib turadi va kompleks tarmoq muhitlarida real vaqt monitoringini ta'minlaydi.

Tadqiqot natijalari asosida aytish mumkinki, neyron tarmoqlar asosidagi tizimlar an'anaviy imzo-asosli usullarga nisbatan tezroq, ishonchliroq va moslashuvchanroq himoya darajasini yaratadi. Shu bilan birga, amaliy joriy etishda ma'lumotlar sifati, modelning resurs talablari va xavfsizlik-maxfiylik masalalariga jiddiy e'tibor qaratish lozim.

Kelajakda ushbu yo'nalishni yanada rivojlantirish uchun federativ o'rganish, Edge-AI va kiberxavfsizlikka yo'naltirilgan o'zbek tilidagi ma'lumotlar bazalarini shakllantirish zarur. Bu esa milliy axborot infratuzilmasida mustaqil, sun'iy intellektga asoslangan kiberxavfsizlik ekotizimini yaratishga xizmat qiladi.

Foydalanilgan adabiyotlar

1. Goodfellow I., Bengio Y., Courville A. *Deep Learning*. MIT Press, 2016.
2. Sahu R., Jena S.K. *AI-Based Phishing Detection Techniques: A Review*. *Journal of Information Security*, 2021.
3. Zhang J. et al. *CNN-LSTM Hybrid Models for Malware Detection*, *IEEE Access*, 2022.
4. O'zbekiston Respublikasi Axborot texnologiyalari va kommunikatsiyalarini rivojlantirish vazirligi. *Kiberxavfsizlik konsepsiyasi*, 2023.
5. Alshamrani A. *AI-Driven Cybersecurity: Challenges and Trends*, *Computers & Security*, 2024.