



Tarmoq hujumlarini aniqlashda sun'iy intellekt texnologiyalarining qo'llanilishi

Axmadaliyev Akramjon Rashidovich

Annotatsiya

Mazkur maqolada zamonaviy tarmoq xavfsizligi muhitida sun'iy intellekt asosidagi yondashuvlarning tarmoq hujumlarini aniqlashdagi o'rni va samaradorligi tahlil qilindi. Mashinaviy o'rganish hamda chuqur o'rganish modellari signaturali va anomaliyaga asoslangan tizimlarga qanday integratsiyalashayotgani, ularning kuchli va zaif tomonlari solishtirildi. O'rganilgan ilmiy manbalar shuni ko'rsatadiki, gibrid arxitekturalar — ya'ni qoidaviy tahlil va o'z-o'zini o'rganuvchi modellar kombinatsiyasi — amaliy muhitda yuqori aniqlik va barqarorlikni ta'minlay oladi. Shuningdek, grafik neyron tarmoqlari (GNN) va tushuntiriluvchan sun'iy intellekt (XAI) texnologiyalari kelgusida kiberxavfsizlik sohasida muhim yo'nalish sifatida shakllanishi mumkinligi asoslab berildi.

Kalit so'zlar: tarmoq hujumlari, sun'iy intellekt, mashinaviy o'rganish, chuqur o'rganish, anomaliya aniqlash, signaturali tahlil, IDS, GNN.

KIRISH

Raqamli transformatsiya jarayonining jadallashuvi tarmoq infratuzilmasini davlat boshqaruvi, moliyaviy tizimlar, sanoat jarayonlari va kundalik hayotning ajralmas qismiga aylantirdi. Shu bilan birga, kiberxavfsizlik masalalari strategik ahamiyat kasb etmoqda. IoT qurilmalari, bulutli xizmatlar va tarqatilgan tizimlarning kengayishi hujum sirtini (attack surface) sezilarli darajada oshirdi. Natijada, tarmoq hujumlari — ruxsatsiz kirish, ma'lumotlarni buzish, xizmatni rad etish (DDoS) yoki yashirin infiltratsiya kabi — murakkablashib bormoqda.

An'anaviy Intrusion Detection Systems (IDS) ko'pincha oldindan aniqlangan imzolar yoki qoidalarga asoslanadi. Masalan, Snort va Suricata kabi tizimlar ma'lum hujum shablonlarini tez aniqlay oladi. Biroq, bunday yondashuvlar ilgari uchramagan ("zero-day") tahdidlar oldida zaif bo'lib qoladi. Bundan tashqari, trafikni shifrlash, tunnellash va obfuskatsiya texnikalari signaturali mexanizmlarni chetlab o'tishga imkon beradi.

So'nggi yillarda sun'iy intellekt, xususan mashinaviy va chuqur o'rganish algoritmlari, ushbu cheklovlarni bartaraf etish vositasi sifatida ko'rilmogda. Ular tarmoq trafigidagi murakkab statistik va vaqtinchalik bog'liqliklarni o'rganib, normal faoliyatdan og'ishlarni aniqlash imkonini beradi. Shunga qaramay, real muhitda AI

asosidagi tizimlarni joriy etishda hisoblash resurslari, noto'g'ri ijobiy natijalar (false positive) va tushuntiriluvchanlik kabi masalalar hanuz dolzarbligicha qolmoqda.

Mazkur ishning maqsadi — sun'iy intellekt texnologiyalarining tarmoq hujumlarini aniqlashdagi rolini tizimli ravishda ko'rib chiqish, mavjud metodlarni taqqoslash hamda istiqbolli yo'nalishlarni aniqlashdir.

TADQIQOT METODOLOGIYASI

Tadqiqot bir necha bosqichda amalga oshirildi. Birinchi bosqichda 2015–2024-yillar oralig'ida chop etilgan ilmiy maqolalar IEEE Xplore, Scopus, Web of Science va Google Scholar bazalaridan tanlab olindi. Tanlash mezonlari sifatida nufuzli jurnallar, yuqori iqtibos ko'rsatkichi va amaliy tajribaga ega ishlar e'tiborga olindi.

Ikkinchi bosqichda maqolalar quyidagi mezonlar bo'yicha tahlil qilindi:

- qo'llanilgan algoritmi turi (supervayzlangan, unsupervayzlangan, reinforcement learning);
- ishlatilgan dataset (KDD Cup 99, NSL-KDD, CICIDS2017, UNSW-NB15 va boshqalar);
- baholash metrikalari (accuracy, precision, recall, F1-score);
- real muhitga tatbiq etish imkoniyati.

Uchinchi bosqichda turli yondashuvlar samaradorlik, moslashuvchanlik va resurs talabi nuqtai nazaridan solishtirildi. Ayniqsa, zero-day hujumlarni aniqlash, ma'lumotlar balanssizligi va modelning umumlashuv qobiliyatiga e'tibor qaratildi.

ADABIYOTLAR SHARHI

Dastlabki IDS tizimlari asosan signaturali tahlilga asoslangan edi. Sommer va Paxson (2010) mashinaviy o'rganish usullarini yopiq (closed-world) modelga asoslangan xavfsizlik muhitida qo'llash cheklovlarini ko'rsatganlar.

Keyinchalik chuqur o'rganish modellarining paydo bo'lishi vaziyatni sezilarli o'zgartirdi. Vinayakumar va boshqalar (2019) DNN asosida yuqori aniqlik ko'rsatkichlariga erishgan. Yin va hamkorlari (2017) esa RNN va LSTM arxitekturalari yordamida vaqt bo'yicha bog'liqliklarni hisobga olish hujumlarni erta bosqichda aniqlash imkonini berishini isbotladilar.

Ferrag va boshqalar (2020) chuqur o'rganish modellarini keng qamrovli taqqoslab, ularning samaradorligi bilan birga hisoblash murakkabligi va ma'lumotlar yetishmovchiligi muammosini ham ta'kidladilar. Khraisat (2019) esa anomaliya asosidagi tizimlarda false positive darajasi yuqori bo'lishi amaliy qo'llashni qiyinlashtirishini ko'rsatgan.

Adabiyotlar tahlili shuni ko'rsatadiki, chuqur o'rganish murakkab naqshlarni aniqlashda samarali bo'lsa-da, umumlashuv va izohlanuvchanlik masalalari hanuz dolzarb.

NATIJALAR VA TAHLIL

AI asosidagi IDS tizimlari ikki asosiy paradigma doirasida rivojlanmoqda: signaturali va anomaliya asosidagi.

Signaturali yondashuv. Random Forest, SVM yoki XGBoost kabi supervayzlangan algoritmlar ma'lum hujum turlarini yuqori aniqlik bilan tasniflay oladi. Ular tez ishlaydi va noto'g'ri aniqlash darajasi nisbatan past. Biroq, yangi va noma'lum hujumlar uchun ular samaradorligini yo'qotadi.

Anomaliya asosidagi yondashuv. Bu tizimlar normal trafik modelini o'rganadi va og'ishlarni aniqlaydi. K-means va DBSCAN kabi klasterlash algoritmlari, Autoencoderlar va LSTM modellar keng qo'llaniladi. Autoencoder normal trafikni qayta tiklaydi, rekonstruksiya xatoligi oshganda anomaliya aniqlanadi. LSTM esa vaqt ketma-ketligidagi dinamikani hisobga oladi.

Graf neyron tarmoqlari (GNN) tarmoqni graf sifatida modellashtirib, tarqatilgan va lateral hujumlarni aniqlashda istiqbolli natijalar bermoqda.

Gibrid arxitekturalar. Amaliy tajribalar shuni ko'rsatadiki, gibrid tizimlar eng barqaror natijani beradi. Signaturali mexanizm tezkor filtr vazifasini bajaradi, anomaliya modeli esa murakkab yoki yangi tahdidlarni aniqlaydi. Ensemble learning usullari noto'g'ri tasniflash ehtimolini kamaytiradi.

Shu bilan birga, real muhitda quyidagi muammolar saqlanib qolmoqda:

- sinflar balanssizligi;
- modelning bir datasetdan boshqasiga yomon umumlashuvi;
- hisoblash resurslariga yuqori talab;
- qarorlarni izohlash qiyinligi.

AI asosidagi IDS tizimlarini samaradorlik, umumlashuv qobiliyati va real muhitga moslashuv darajasi bo'yicha solishtirish maqsadida turli yondashuvlar tahlil qilindi. Quyidagi jadval asosiy paradigmlar o'rtasidagi farqlarni umumlashtiradi.

1-jadval. IDS yondashuvlarining taqqoslovchi tahlili

Yondashuv turi	Qo'llaniladigan algoritmlar	Zero-day aniqlash	O'rtacha aniqlik darajasi	Real vaqtga moslik	Hisoblash talabi
Signaturali	Snort, Suricata (qoidaviy)	Past	80–90%	Juda yuqori	Past
Supervayzlangan ML	RF, SVM, XGBoost	O'rta	90–97%	Yuqori	O'rta
Klasterlash (unsupervised)	K-means, DBSCAN	O'rta-yuqori	85–93%	O'rta	O'rta
Autoencoder	Deep Autoencoder	Yuqori	90–96%	O'rta	Yuqori
RNN/LSTM	LSTM, GRU	Yuqori	93–98%	O'rta	Juda yuqori
GNN	Graph Convolutional Network	Juda yuqori (lateral hujumlarda)	92–97%	Past-o'rta	Juda yuqori
Gibrid tizimlar	RF + Autoencoder, IDS + DL	Eng yuqori	95–99%	Yuqori	Yuqori

Jadvaldan ko'rinadiki, an'anaviy signaturali tizimlar real vaqt rejimida ishlashda ustunlikka ega bo'lsa-da, zero-day hujumlarni aniqlashda zaifdir. Supervayzlangan

mashinaviy o'rganish algoritmlari yuqori aniqlik beradi, biroq ular uchun katta hajmdagi belgilangan (labeled) ma'lumotlar talab etiladi.

Chuqur o'rganish modellari, xususan LSTM va Autoencoder arxitekturalari, murakkab va vaqtga bog'liq naqshlarni aniqlashda yuqori natija ko'rsatadi. Biroq, ular hisoblash resurslariga talabchan bo'lib, real vaqt sharoitida optimallashtirishni talab qiladi.

Graf neyron tarmoqlari ayniqsa lateral harakatlanish (lateral movement) va tarqatilgan hujumlarni aniqlashda istiqbolli yo'nalish hisoblanadi, chunki ular tarmoq topologiyasini hisobga oladi. Ammo ularning implementatsiyasi murakkab va hisoblash jihatdan og'ir.

Eng barqaror natijalar gibrid arxitekturalarda kuzatiladi. Bunday tizimlar an'anaviy tezkor filtr mexanizmi bilan chuqur o'rganish modelini birlashtiradi, bu esa aniqlik va moslashuvchanlikni bir vaqtning o'zida ta'minlaydi.

XULOSA VA ISTIQBOLLAR

Tahlil natijalari sun'iy intellekt asosidagi yondashuvlar tarmoq hujumlarini aniqlashda sezilarli ustunlikka ega ekanini ko'rsatadi. Ayniqsa, vaqtga bog'liq va yashirin naqshlarni aniqlashda chuqur o'rganish modellari samarali.

Biroq, AI tizimlari universal yechim emas. Ularning samaradorligi ma'lumot sifati, xususiyatlar tanlovi va implementatsiya strategiyasiga bog'liq. Shuningdek, real vaqt rejimida ishlash va tushuntiriluvchanlik muammolari dolzarb bo'lib qolmoqda.

Kelgusida quyidagi yo'nalishlar istiqbolli deb hisoblanadi:

- GNN va Transformer modellarini tarmoq oqimlarini tahlil qilishga moslashtirish;
- XAI vositalarini IDS tizimlariga integratsiya qilish;
- model siqish va optimallashtirish orqali real vaqt samaradorligini oshirish;
- adversarial training orqali zero-day hujumlarga chidamlilikni kuchaytirish.

Shunday qilib, sun'iy intellekt texnologiyalari tarmoq xavfsizligida muhim o'rin egallamoqda. Biroq, ularni samarali joriy etish texnik yechimlar bilan bir qatorda, amaliy tajriba va soha mutaxassislari hamkorligini ham talab etadi.

Foydalanilgan adabiyotlar

[1]. Sommer R., Paxson V. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection // IEEE Security & Privacy. – 2010. – Vol. 8, № 6. – P. 38–46.

[2]. Vinayakumar R., Alazab M., Soman K. P., Poornachandran P. Deep Learning Approach for Intelligent Intrusion Detection System // IEEE Access. – 2019. – Vol. 7. – P. 41525–41550.

[3]. Yin C., Zhu Y., Fei J., He X. A Deep Learning Approach for Intrusion Detection Using Recurrent Neural Networks // IEEE Access. – 2017. – Vol. 5. – P. 21954–21961.

[4]. Ferrag M. A., Maglaras L., Moschoyiannis S., Janicke H. Deep Learning for Cyber Security Intrusion Detection: Approaches, Datasets, and Comparative Study // Journal of Information Security and Applications. – 2020. – Vol. 50. – Article 102419.

[5]. Khraisat A., Gondal I., Vamplew P., Kamruzzaman J. Survey of Intrusion Detection Systems: Techniques, Datasets and Challenges // Security and Communication Networks. – 2019. – Vol. 2019. – Article ID 3402858.

[6]. Sarhan M., Layeghy S., Portmann M. Evaluating Standard Feature Sets towards Increased Generalisability and Explainability of ML-Based Network Intrusion Detection // Journal of Information Security and Applications. – 2021. – Vol. 63. – Article 103045.

[7]. Chawla N. V., Bowyer K. W., Hall L. O., Kegelmeyer W. P. SMOTE: Synthetic Minority Over-sampling Technique // Journal of Artificial Intelligence Research. – 2002. – Vol. 16. – P. 321–357.