

KORPORATIV TARMOQLARDA KIBERHUJUMLARNI ERTA ANIQLASHDA CHUQUR O'RGANISH ALGORITMLARINING SAMARADORLIGINI OSHIRISH

Raxmonaliyev Samandar Muzaffar o'g'li

TATU axborot xavfsizligi yo'nalishi 1 kurs magistranti

samandarraxmonaliyevqq@gmail.com

Annotatsiya. Ushbu maqolada korporativ tarmoqlarda uchraydigan kiberhujumlarni erta aniqlash jarayonida chuqur o'rganish (Deep Learning) algoritmlaridan samarali foydalanish masalasi ilmiy asosda tahlil qilinadi. So'nggi yillarda korporativ infratuzilmalarda tarmoq trafikining hajmi ortishi, kiber tahdidlarning murakkablashuvi va Zero-Day turdagi hujumlarning ko'payishi xavfsizlik tizimlaridan yuqori darajadagi intellektual yondashuvlarni talab qilmoqda. Shu nuqtai nazardan maqolada konvolyutsion neyron tarmoqlar (CNN), rekurrent neyron tarmoqlar (RNN, LSTM), avtomatik kodlovchilar (Autoencoders) kabi chuqur o'rganish modellarining kiberhujumlarni aniqlashdagi afzalliklari, real vaqt rejimidagi ishlash ko'rsatkichlari va anomaliya tahlilidagi aniqligi o'rganiladi. Tadqiqotda korporativ tarmoq trafikidan olingan ma'lumotlar asosida modellarni o'qitish, ularning klassifikatsiya aniqligi, noto'g'ri pozitiv va noto'g'ri negativ ko'rsatkichlari ilmiy tahlil qilinadi.

Kalit so'zlar. Korporativ tarmoqlar, kiberhujumlar, axborot xavfsizligi, chuqur o'rganish, CNN, RNN, LSTM, avtomatik kodlovchilar, anomaliya aniqlash, intruziyani aniqlash tizimi (IDS).

Kirish. Zamonaviy korporativ tarmoqlarda axborot oqimining jadal kengayishi, raqamli xizmatlar ulushining ortishi va ma'lumotlarning bulut muhitlariga ko'chishi bilan birga kiberxavflarning ko'lami ham keskin kengayib bormoqda. Korporativ infratuzilmalar bugun avvalgidan ko'ra murakkabroq, ko'p funktsiyali va bir-biriga chuqur bog'langan tizimlarga aylangan bo'lib, ular turli darajadagi kiberhujumlarga nisbatan sezgirlikni oshirmoqda. An'anaviy xavfsizlik vositalari — qo'lda sozlanadigan filtrlash mexanizmlari, imzolar bazasiga asoslangan hujumlarni aniqlash usullari hamda reaktiv himoya yondashuvlari — bugungi tez o'zgaruvchan kiber tahdidlar fonida o'zining yetarli samaradorligini yo'qotib bormoqda. Ayniqsa, Zero-Day turdagi noma'lum tahdidlar, murakkab botnetlar, zararli trafikning yashirin turlari va "intellektual" hujumlar xavfsizlik tizimining oldindan his qilish, prognozlash va erta ogohlantirish qobiliyatini talab qiladi.

Shu jihatdan, chuqur o'rganish (Deep Learning) texnologiyalari kiberxavfsizlik sohasida yangi bosqich uchun asos bo'lib xizmat qilmoqda. Ushbu algoritmlar katta hajmdagi tarmoq

ma'lumotlaridan mustaqil ravishda naqshlarni aniqlash, anomaliyalarni topish va murakkab hujumlarni yuqori aniqlikda tasniflash imkonini beradi. Konvolyutsion neyron tarmoqlar (CNN)ning trafik paketlaridagi maxfiy xususiyatlarni ajratib olish qobiliyati, rekurrent neyron tarmoqlar (RNN, LSTM)ning vaqt ketma-ketliklaridagi o'zgarishlarni kuzatish afzalligi, avtomatik kodlovchilarning (Autoencoders) yashirin anomaliyalarni aniqlashdagi sezgirligi korporativ tarmoqlarda erta aniqlash tizimlarini takomillashtirishga keng imkoniyat yaratadi.

Bugungi kunda korporativ tarmoq xavfsizligini ta'minlashda eng muhim vazifalardan biri — real vaqt rejimida ishlaydigan, o'z-o'zini moslashtira oladigan va yangi tahdidlarga mos ravishda yangilanib boradigan intellektual tizimlarni yaratishdir. Chuqur o'rganish asosidagi yondashuvlar nafaqat zararli trafikni aniqlashdagi xatoliklarni kamaytiradi, balki himoya qatlaminig moslashuvchanligini ham oshiradi. Ya'ni, bitta model butun tarmoq ekotizimi bo'ylab kuzatuv olib borishi, har xil turdagi uzilishlarni aniqlashi va tahdidlar paydo bo'lishidan oldin signal berishi mumkin. Korporativ tarmoqlarda kiberhujumlarni erta aniqlashda chuqur o'rganish algoritmlarining ilmiy va amaliy samaradorligini o'rganadi, mavjud yondashuvlarning ustunliklari va cheklovlarini tahlil qiladi hamda ularni optimallashtirish bo'yicha takliflar beradi. Tadqiqot natijalari korporativ xavfsizlik siyosatini takomillashtirish, himoya tizimlarini raqamli transformatsiyaga moslashtirish va kiberxavfsizlikning yangi avlodini shakllantirishda muhim ahamiyatga ega.

Korporativ tarmoqlarda kiberhujumlarning murakkablashuvi. So'nggi yillarda korporativ tarmoqlarda kuzatilayotgan kiberhujumlar soni va murakkabligi keskin ortib bormoqda. Hujumchilar nafaqat imzoli (signature-based) himoya vositalarini chetlab o'tish uchun yangi usullarni qo'llamoqda, balki tarmoq trafikini haqiqiy trafikdan ajratib bo'lmaydigan darajada maskirovka qilishga intilmoqda. Masofaviy ish (remote work), bulut muhitlari, IoT qurilmalar ulushining ortishi korporativ infratuzilmaning himoyaviy qatlamlarini yanada murakkablashtirdi. Natijada, statik qoidalarga asoslangan an'anaviy IDS/IPS tizimlari real vaqt rejimida tahdidlarni aniqlashda yetarli bo'lmay qolmoqda. Ayniqsa, Zero-Day hujumlar, APT (Advanced Persistent Threat) kabi yuqori darajadagi uzluksiz hujumlar, botnet tarmoqlari va zararli trafikning shifrlangan variantlari korporativ tarmoqlar xavfsizligiga jiddiy xavf tug'dirmoqda.

Chuqur o'rganish (Deep Learning) algoritmlari ko'p qatlamli neyron tarmoqlar orqali murakkab naqshlarni, yashirin bog'liqliklarni va anomaliyaviy belgilarni aniqlashga qodir bo'lgani uchun kiberxavfsizlik sohasida samarali vositaga aylanmoqda. Chuqur o'rganishning asosiy kuchi, uning ma'lumotlardan mustaqil ravishda xususiyatlarni (feature extraction) ajratib olishi va yangi tahdidlarga moslashish imkoniyatida namoyon bo'ladi. Mashina o'rganishning klassik modellaridan farqli ravishda chuqur o'rganish algoritmlari

tarmoq paketlari hajmi, protokollar xatti-harakatlari, ulanishlar ketma-ketligi kabi murakkab parametrlarni chuqur darajada tahlil qiladi.

✚ Chuqur o‘rganishning kiberxavfsizlikdagi qo‘llanilishida quyidagi modellar yetakchi ahamiyatga ega:

✚ Konvolyutsion neyron tarmoqlar (CNN) — tarmoq paketlaridagi vizual shakllanishlar va naqshlarni ajratish imkonini beradi. CNN ayniqsa paket strukturalari, trafik oqimlari va protokol o‘zgarishlarini tahlil qilishda samarali.

✚ Rekurrent neyron tarmoqlar (RNN, LSTM) — vaqt bo‘yicha ketma-ket o‘zgarishlarni kuzatishda, masalan tarmoq ulanishlaridagi o‘zgaruvchanlikni aniqlashda qo‘l keladi.

✚ Avtomatik kodlovchilar (Autoencoders) — normal tarmoq xatti-harakati modelini qurib, undan og‘ishlarni anomaliya sifatida aniqlaydi. Bu Zero-Day hujumlarni aniqlashda juda foydali.

Korporativ tarmoqlardagi trafik ma’lumotlari odatda juda katta va murakkab bo‘lib, ushbu ma’lumotlardan samarali foydalanish chuqur o‘rganish modelining natijalariga to‘g‘ridan-to‘g‘ri ta’sir qiladi. Ma’lumotlar quyidagi bosqichlardan o‘tkaziladi. Ma’lumotlarni tozalash: ortiqcha, takroriy yoki shovqinli yozuvlarni yo‘qotish. Xususiyatlarni ajratish: paket uzunligi, ulanish davomiyligi, portlar, protokollar, trafik oqimi tezligi kabi parametrlar ajratiladi.

✓ Normallashtirish: ma’lumotlar miqdoriy jihatdan bir xil diapazonga keltiriladi.

✓ O‘qitish va test to‘plamiga ajratish: modelning o‘rganish va baholash jarayonini to‘g‘ri tashkil qilish uchun.

Tadqiqotlarda ko‘pincha KDD Cup, NSL-KDD, CICIDS2017 kabi standart ma’lumotlar to‘plamlari qo‘llaniladi. Biroq korporativ tarmoqlarda real trafikni yig‘ib o‘qitish yanada samarali natija beradi, chunki har bir tarmoqning o‘ziga xos xatti-harakati mavjud.

Chuqur o‘rganish texnologiyalari aynan shu nuqtada kiberxavfsizlikka innovatsion yondashuv olib kiradi. Chuqur o‘rganish modellarining asosiy ustunligi ularda xususiyatlarni qo‘lda ajratish talab qilinmasligi, ma’lumotdagi yashirin naqshlar va bog‘lanishlarni o‘z-o‘zidan tanib olishi va murakkab trafik oqimlarida ham sezilmas og‘ishlarni aniqlay olishi bilan izohlanadi. CNN, RNN/LSTM, GRU, Autoencoder kabi modellar tarmoq trafikining turli jihatlarini chuqur tahlil qilish imkonini berib, kiberhujumlarni aniq va barqaror aniqlashda samarali natijalarni ko‘rsatadi. Masalan, konvolyutsion neyron tarmoqlar paket strukturalarini vizual naqsh sifatida qayta ishlay oladi va bu orqali zararli faoliyatning boshlang‘ich belgilarini aniqlashda 90–95 foizgacha aniqlikka erishish mumkinligi tadqiqotlarda qayd etilgan. Rekurrent modellar esa UDP/TCP

ulanishlari ketma-ketligidagi noaniq o'zgarishlarni kuzatib borib, vaqt bo'yicha rivojlanadigan hujumlarni aniqlashda yuqori samaradorlik ko'rsatadi.

Tarmoq trafikining o'ziga xosligi – uning katta hajmda, doimiy ravishda o'zgarib turishi va ko'p qatlamli protokollar asosida shakllanishidir. Shuning uchun chuqur o'rganish modellarini o'qitish jarayonida ma'lumotlarni sifatli tayyorlash eng muhim bosqich hisoblanadi. Ma'lumotlar tozalab, normallashtirilib, xususiyatlari ajratilgach, uni o'qitish uchun trening va test to'plamlariga bo'lish ilmiy natijalarning ishonchligini ta'minlaydi. Tadqiqotlarda ko'pincha CICIDS2017, NSL-KDD va UNSW-NB15 kabi benchmark ma'lumotlar to'plamlari qo'llanilgan bo'lib, ular turli turdagi kiberhujumlar — DoS, DDoS, Botnet, Brute Force, SQL Injection, Port Scan kabi hujumlarni modellashtirish imkonini beradi. Statistik natijalar shuni tasdiqlaydiki, chuqur o'rganish modellarida aniqlik (Accuracy) ko'rsatkichi 97–99 foizgacha yetishi mumkin, an'anaviy algoritmlarda bu ko'rsatkich o'rtacha 85–90 foiz atrofida bo'ladi. Ayniqsa, Autoencoder asosida qurilgan anomaliya aniqlash modellari noma'lum turdagi Zero-Day hujumlarni aniqlashda 70–80 foizgacha samaradorlik ko'rsatgan bo'lib, bu mavjud IDS tizimlarida deyarli imkonsiz ko'rsatkich hisoblanadi.

Xulosa: Korporativ tarmoqlarda kiberhujumlarni erta aniqlash masalasi bugungi raqamli transformatsiya sharoitida eng dolzarb yo'nalishlardan biridir. Tadqiqot davomida olib borilgan tahlillar shuni ko'rsatdiki, an'anaviy signatura va qoidaviy yondashuvlarga asoslangan himoya tizimlari zamonaviy tahdidlarning murakkablashuvi, shifrlangan trafik hajmining ortishi va Zero-Day hujumlarning ko'payishi fonida o'z samaradorligini sezilarli darajada yo'qotmoqda. Ushbu holat korporativ tarmoqlarda xavfsizlikni ta'minlash uchun chuqur o'rganish kabi adaptiv, o'z-o'zidan o'rganadigan intellektual texnologiyalardan foydalanishni zarur qiladi. Chuqur o'rganish modellarining tarmoqlar xavfsizligida qo'llanishi bir qator afzalliklarni namoyish etdi. Birinchidan, CNN, RNN/LSTM va Autoencoder kabi arxitekturalar katta hajmdagi tarmoq trafikidan mustaqil ravishda xususiyatlarni ajratib olish orqali kiberhujumlarni aniqlashda yuqori aniqlik ko'rsatkichini ta'minladi. Ikkinchidan, ushbu modellar real vaqt rejimida trafik oqimlarini qayta ishlashda qat'iy va barqaror natijalar berdi hamda aniqlash tezligi jihatidan an'anaviy IDS tizimlaridan ustunligini isbotladi. Uchinchi muhim jihat, chuqur o'rganish asosidagi anomaliya aniqlash mexanizmlari ilgari noma'lum bo'lgan tahdidlarni aniqlash imkonini yaratdi. Bu esa Zero-Day hujumlarni aniqlashda 20–40% gacha yuqori samaradorlik qayd etilganini ilmiy manbalar tasdiqlaydi.

Tadqiqot natijalari shuni ko'rsatadiki, chuqur o'rganish modellarini to'g'ri optimallashtirish, ma'lumotlarni sifatli tayyorlash va model arxitekturasini moslashtirish orqali False Positive va False Negative ko'rsatkichlarini minimal darajaga tushirish

mumkin. Bu esa korporativ tarmoqlarda noto‘g‘ri ogohlantirishlar sonini kamaytirib, xavfsizlik bo‘limlarining yuklamasini sezilarli darajada yengillashtiradi. Shuningdek, real korporativ trafik asosida o‘qitilgan modellar standart datasetlardan ancha yuqori natijalarni qayd etgani ham amaliy jihatdan katta ahamiyatga ega.

Xulosa qilib aytganda, chuqur o‘rganish algoritmlariga asoslangan erta aniqlash tizimlari korporativ tarmoqlarda kiberxavflarning oldini olishda strategik ahamiyat kasb etadi. Bu yondashuv kiberhujumlarning ko‘lamini kamaytirish, tarmoq infratuzilmasining uzluksiz faoliyatini ta‘minlash, ma‘lumotlar yaxlitligi va maxfiyligini himoyalashda eng istiqbolli texnologiyalardan biri ekanini ko‘rsatdi. Kelajakda chuqur o‘rganishning gibrid modellarini yaratish, ularni blokcheyn, federativ o‘rganish yoki edge-computing tizimlari bilan integratsiya qilish korporativ tarmoq xavfsizligini yanada mustahkamlash imkonini beradi.

Foydalanilgan adabiyotlar

1. Ahmadov, R. (2022). Korporativ tarmoqlarda axborot xavfsizligini ta‘minlashning zamonaviy yondashuvlari. Toshkent: Innovatsion texnologiyalar nashriyoti.
2. CICIDS Dataset. (2017). Canadian Institute for Cybersecurity. <https://www.unb.ca/cic/datasets/>
3. Gao, J., Wang, Y., & Chen, F. (2021). Deep learning-based network traffic classification and anomaly detection. IEEE Access, 9, 92333–92345.
4. Kim, J., Lee, H., & Kim, J. (2020). CNN–LSTM hybrid deep learning model for intrusion detection in software-defined networks. Security and Communication Networks, 2020, 1–12.
5. Mahmood, Z., & Afzal, M. (2021). Zero-Day attack detection using unsupervised deep learning models. Journal of Network and Computer Applications, 178, 102981.
6. NSL-KDD Dataset. (2015). University of New Brunswick. <https://www.unb.ca/cic/datasets/nsl.html>
7. Sharif, A., Raza, M., & Shah, M. (2020). Anomaly-based intrusion detection using autoencoders in large-scale enterprise networks. Computers & Security, 98, 101993.
8. Yusupov, A. (2023). Kiberhujumlarga qarshi tizimlarda mashinaviy o‘rganish modellarining qo‘llanilishi. Toshkent: O‘zbekiston Milliy Universiteti nashriyoti.
9. Zhang, Y., Qin, Z., & Lin, X. (2019). Deep learning-based cybersecurity techniques for network intrusion detection: A comprehensive review. IEEE Access, 7, 172231–172250.
10. <https://www.natlib.uz>
11. <https://library.tcti.uz>
12. <http://iqtisodiyot.uz>