

KIBERXAVFSIZLIK VA SUNIY INTELLEKT INTEGRATSIYASI

Toshkent Amaliy Fanlar Universiteti

“Tarmoqlar iqtisodiyoti va buhgalteriyasi” kafedrası assistenti

Xaqqulova Ra’no Abduganiyevna

Toshkent Amaliy Fanlar Universiteti ECO2301 talabasi

Atajonov Umidjon Shamrat o’g’li

Annotatsiya: Ushbu maqolada kiberxavfsizlikni ta’minlashda sun’iy intellekt (SI) texnologiyalarini integratsiya qilish masalalari tadqiq etiladi. Zamonaviy kiberhujumlarning murakkablashishi sharoitida mashinali o’rganish va chuqur o’rganish algoritmlarining tahdidlarni aniqlash va ularga qarshi harakat qilishdagi samaradorligi tahlil qilingan. Shuningdek, maqolada SI tizimlaridan ham himoya, ham hujum vositasi sifatida foydalanishning ikki yoqlama xususiyati va kelajakdagi xavfsizlik arxitekturasini yaratish istiqbollari yoritilgan.

Аннотация: В данной статье исследуются вопросы интеграции технологий искусственного интеллекта (ИИ) в сферу кибербезопасности. В условиях усложнения современных кибератак анализируется эффективность алгоритмов машинного и глубокого обучения в обнаружении угроз и реагировании на них. Также рассматривается двойственная природа ИИ как инструмента защиты и нападения, а также перспективы создания архитектур безопасности будущего на основе интеллектуальных систем.

Abstract: This article explores the integration of Artificial Intelligence (AI) technologies in ensuring cybersecurity. In the context of increasingly complex cyber threats, the effectiveness of machine learning and deep learning algorithms in threat detection and response is analyzed. The paper also discusses the dual nature of AI as both a defensive and offensive tool, as well as the prospects for developing future security architectures based on intelligent systems.

Kalit so’zlar: kiberxavfsizlik, sun’iy intellekt, mashinali o’rganish, anomaliyalarni aniqlash, kiberhujum, ma’lumotlar himoyasi.

Ключевые слова: кибербезопасность, искусственный интеллект, машинное обучение, обнаружение аномалий, кибератака, защита данных.

Keywords: cybersecurity, artificial intelligence, machine learning, anomaly detection, cyberattack, data protection.

I. Kirish:

Mavzuning dolzarbligi va maqsadlari

Bugungi kunda jahon iqtisodiyoti va ijtimoiy hayotining jadallik bilan raqamlashishi nafaqat yangi imkoniyatlar, balki misli ko’rilmagan kiber-tahdidlarni ham keltirib chiqarmoqda. An’anaviy kiberhujumlar o’rnini endilikda avtomatlashtirilgan, o’zgaruvchan va yuqori darajada nishonli hujumlar egallamoqda. Shu sababli, axborot tizimlarini himoya qilishda faqat inson resurslariga va qat’iy qoidalarga (rule-based systems) tayangan klassik usullar o’z samaradorligini yo’qotmoqda. Ushbu nuqtada Sun’iy intellekt (SI) va kiberxavfsizlik integratsiyasi zamonaviy mudofaa strategiyasining markaziy bo’g’iniga aylandi.

Zamonaviy kiberhujumlar soni sekundiga bir necha mingdan oshib ketishi mumkin. Inson-operator bunday hajmdagi ma’lumotlarni tahlil qilishga va real vaqt rejimida qaror qabul qilishga jismonan ulgurmaydi. SI algoritmlari esa millisekundlar ichida anomal holatlarni aniqlash va ularga chek qo’yish qobiliyatiga ega.

Eskicha antivirus va xavfsizlik devorlari faqat ma’lum bo’lgan viruslar bazasi asosida ishlaydi. Biroq, hali dunyoda qayd etilmagan yangi turdagi xavflarni aniqlash uchun SI tarkibidagi Mashinali o’rganish (Machine Learning) modellarining prognozlash xususiyati nihoyatda muhim.

Ular ma'lumotlar bazasiga tayanmasdan, xulq-atvor tahlili (behavioral analysis) orqali noma'lum xavfni "sezish" imkonini beradi.

Dolzarblikning eng xavfli jihati shundaki, tajovuzkorlar ham o'z hujumlarini avtomatlashtirish uchun SI vositalaridan foydalanishmoqda. Murakkab fishing xabarlar, deepfake texnologiyalari va adaptiv zararli dasturlarga qarshi faqatgina ulardan kuchliroq bo'lgan "Aqlli mudofaa" tizimi bilan kurashish mumkin.

Korporativ tarmoqlar har kuni terabaytlab log-fayllarni ishlab chiqaradi. Ushbu ma'lumotlar oqimi ichidan shubhali nuqtani topish xuddi "samo somon ichidan igna izlash"ga o'xshaydi. SI algoritmlari esa bu "somon" ichidagi qonuniyatlarni ajratib ko'rsata oladigan yagona vositadir.

Hozirgi kunda kiberxavfsizlik mutaxassislar yetishmovchiligi global miqyosda sezilmoqda. SI ushbu vakuum o'rnini to'ldirish, muntazam va bir xil zerikarli operatsiyalarni (monitoring, skanerlash) o'z zimmasiga olish va inson kapitalini faqat murakkab strategik qarorlar qabul qilishga yo'naltirish imkonini beradi.

Shunday qilib, kiberxavfsizlikda SI integratsiyasi shunchaki texnologik yangilanish emas, balki raqamli suverenitet va axborot xavfsizligini ta'minlashning fundamental asosi hisoblanadi.

II. Asosiy qism

2.1. Kiberxavfsizlikda SI uslubiyoti: Mashinali o'rganish (ML) va Chuqur o'rganish (DL) modellarining anomaliyalarni aniqlashdagi roli.

Zamonaviy kiber-himoya tizimlarining intellektuallashuvi bevosita Mashinali o'rganish (ML) va Chuqur o'rganish (DL) algoritmlarining muvaffaqiyatli integratsiyasiga bog'liq. Ushbu texnologiyalar kiberxavfsizlikni "reaktiv" (hujumdan keyin chora ko'rish) holatidan "proaktiv" (hujumni sodir bo'lmasidan avval aniqlash) holatiga o'tkazish imkonini beradi.

ML modellarining anomaliyalarni aniqlashdagi asosiy roli — tizim va tarmoq trafigining "normal" holatini matematik model sifatida shakllantirishdir. Bunda asosan ikki xil o'qitish usuli qo'llaniladi.

Nazorat ostidagi o'qitish (Supervised Learning): Bu usulda algoritm ma'lum bir "etiketlangan" (labeled) ma'lumotlar to'plami (masalan, "zararli" va "xavfsiz" trafik namunalari) yordamida o'qitiladi. Support Vector Machines (SVM) va Random Forest kabi algoritmlar ma'lum turdagi hujumlarni (DDoS, SQL injection) yuqori aniqlikda klassifikatsiya qiladi.

Nazoratsiz o'qitish (Unsupervised Learning): Bu kiberxavfsizlikda eng muhim yo'nalishdir, chunki u ilgari noma'lum bo'lgan (Zero-day) hujumlarni aniqlash imkonini beradi. Clustering (klasterlash) usullari orqali algoritm tarmoqdagi noodatiy xatti-harakatlarni guruhlaydi va ularni anomaliya sifatida bayroq ostiga oladi.

DL kiberxavfsizlikka inson aralashuvi minimal bo'lgan, ko'p qatlamli neyron tarmoqlarini olib kirdi. Uning MLdan ustunligi — xususiyatlarni avtomatik ajratib olish (automatic feature extraction) qobiliyatidir.

Recurrent Neural Networks (RNN) va ayniqsa LSTM (Long Short-Term Memory) tarmoqlari vaqtinchalik ketma-ketlikka ega bo'lgan hujumlarni tahlil qilishda samaralidir. Masalan, sekin amalga oshiriladigan (low-and-slow) APT hujumlari uzoq vaqt davomida kichik izlar qoldiradi; DL modellari ushbu tarqoq izlarni yaxlit xavf sifatida birlashtira oladi.

DL yordamida zararli dastur kodi vizuallashadi yoki baynar darajada tahlil qilinadi. Bu esa kiberjinoyatchilar tomonidan qo'llaniladigan kodni chalkashtirish (obfuscation) usullarini aylanib o'tishga yordam beradi.

Anomaliyalarni aniqlashda SI uslubiyotining asosiy vazifasi "False Positive" (noto'g'ri ijobiy natija) ko'rsatkichini kamaytirishdir. An'anaviy tizimlar har qanday noodatiy yuklamani

xavf deb hisoblasa, SI modellari kontekstni tushunadi (masalan, bayram kunlari trafikning ortishi hujum emas, balki tabiiy jarayon ekanini anglaydi).

Xulosa qilib aytganda, ML va DL modellarining integratsiyasi kiber-mudofaani avtomatlashtirilgan, o'z-o'zini o'qitadigan va dinamik o'zgaruvchan ekotizimga aylantiradi. Bu esa soniyalar ichida millionlab operatsiyalarni tahlil qilish va inson ko'zi ilg'amas tahdidlarni bartaraf etishning yagona samarali yo'lidir.

2.2. Amaliy tatbiq sohalari

Sun'iy intellekt kiberxavfsizlikda shunchaki nazariya emas, balki aniq amaliy muammolarni hal qiluvchi vositadir. Quyida uning eng muhim qo'llanilish sohalari keltirilgan.

Tarmoq xavfsizligi: Trafikni real vaqt rejimida tahlil qilish

Tarmoq xavfsizligida SI tizimlari "Smart Firewalls" va Intrusion Detection Systems (IDS) asosini tashkil etadi.

Real-time Analysis: An'anaviy tizimlar ma'lumotlar paketini faqat sarlavhasi (header) bo'yicha tekshirsa, SI asosidagi tizimlar har bir paketning mazmunini va oqimlar orasidagi mantiqiy bog'liqlikni millisekundlar ichida tahlil qiladi.

DDoS hujumlaridan himoya: SI trafingining keskin o'sishini tahlil qilib, uning qonuniy foydalanuvchilar oqimi yoki botnet hujumi ekanligini darhol ajratadi va zararli so'rovlarni avtomatik bloklaydi.

Endpoint Protection: Malwareni xatti-harakati orqali aniqlash

Endpoint (oxirgi nuqta) himoyasi — bu foydalanuvchi qurilmalari (kompyuterlar, serverlar, mobil qurilmalar) darajasidagi xavfsizlikdir.

Behavioral Analysis: Bugungi kunda zararli dasturlar (malware) o'z kodini doimiy ravishda o'zgartirib turadi (polimorfizm). SI bu dasturning "kimligi"ga (fayl nomiga yoki imzosiga) emas, balki "nima qilayotgani"ga e'tibor beradi.

Pre-execution blocking: Masalan, agar oddiy matn muharriri kutilmaganda tizim fayllarini shifrlashni (ransomware kabi) boshlasa, SI algoritmi ushbu jarayonni g'ayritabiiy deb hisoblaydi va jarayonni to'xtatib, foydalanuvchini ogohlantiradi.

Antifishing: NLP yordamida ijtimoiy muhandislikni bloklash

Fishing hujumlarining 90% dan ortig'i inson psixologiyasiga (ijtimoiy muhandislik) asoslangan. NLP (Natural Language Processing) texnologiyalari bu yerda raqamli "psixolog" vazifasini bajaradi.

Semantik tahlil: SI xatning matnini o'qib, undagi "shoshilinchlik" (masalan, "Hisobingiz 1 soat ichida bloklanadi!"), "qo'rqitish" yoki "soxta manfaat" elementlarini aniqlaydi.

Yashirin xavflarni aniqlash: NLP tizimlari vizual jihatdan bir xil ko'ringan, ammo texnik jihatdan farq qiluvchi belgilarni (masalan, google.com o'rniga g00gle.com) va xat jo'natuvchining odatiy uslubidan og'ishlarni aniqlab, xatni xavfli deb belgilaydi.

Muhim jihat: Ushbu amaliy sohalarning barchasida SI insonning xatolarini (charchoq, e'tiborsizlik) kompensatsiya qiladi va tizimning 24/7 rejimida uzluksiz va yuqori aniqlikda ishlashini ta'minlaydi.

2.3. Integratsiyaning xavf-xatarlari: "Adversarial AI" va kiberjinoyatchilar tomonidan SIning qo'llanilishi

Kiberxavfsizlikda sun'iy intellekt nafaqat mudofaa qalqoni, balki hujumkor qurol sifatida ham namoyon bo'lmoqda. SI tizimlarining joriy etilishi yangi turdagi zaifliklarni va "aqli" hujum vektorlarini vujudga keltirdi.

Adversarial Machine Learning (Raqibona mashinali o'rganish)

Adversarial AI — bu SI modellarining oʻzini aldashga yoki ishdan chiqarishga qaratilgan hujum uslubidir. Bunda tajovuzkorlar algoritmlarning qaror qabul qilish mexanizmlaridagi kamchiliklardan foydalanadilar:

Evasion Attacks (Aylanib oʻtish hujumlari): Hujumchi zararli dastur (malware) kodiga inson koʻzi ilgʻamas kichik oʻzgarishlar kiritadi. Bu oʻzgarishlar SI modelini ushbu faylni "toza" deb tasniflashga majbur qiladi.

Data Poisoning (Maʼlumotlarni zaharlash): Agar SI tizimi doimiy ravishda yangi maʼlumotlar asosida oʻzini oʻqitib borsa, tajovuzkorlar oʻquv bazasiga notoʻgʻri maʼlumotlarni kiritish orqali modelning "mantiqiy filtri"ni buzishi mumkin. Natijada, tizim kelajakda haqiqiy hujumlarni oddiy holat deb qabul qila boshlaydi.

Kiberjinoyatchilar qoʻlida SI: Hujumlarni avtomatlashtirish

Kiberjinoyatchilar hujum samaradorligini oshirish va xarajatlarni kamaytirish uchun SI imkoniyatlaridan keng foydalanmoqda:

AI-Enhanced Phishing: SI yordamida yaratilgan fishing xatlari inson tomonidan yozilganidan deyarli farq qilmaydi. Tabiiy tilni qayta ishlash (NLP) modellari orqali minglab foydalanuvchilarning shaxsiy uslubiga moslashtirilgan (spear-phishing) xabarlarini avtomatik generatsiya qilish imkoniyati paydo boʻldi.

Deepfakes va Ijtimoiy muhandislik: Ovozli va video deepfake texnologiyalari orqali kompaniya rahbarlarining qiyofasini soxtalashtirish, moliyaviy oʻtkazmalarni amalga oshirish yoki maxfiy maʼlumotlarni qoʻlga kiritish holatlari koʻpaymoqda.

Adaptiv zararli dasturlar: SIDan foydalanuvchi viruslar tarmoq ichidagi xavfsizlik choralarini mustaqil tahlil qila oladi va aniqlanmaslik uchun oʻz xatti-harakatlarini dinamik ravishda oʻzgartiradi.

Algoritmik ishonchsizlik va "Qora quti" muammosi

Koʻplab SI modellari (ayniqsa chuqur oʻrganish neyron tarmoqlari) "qora quti" (black box) tamoyili asosida ishlaydi. Yaʼni, tizim nima uchun maʼlum bir trafikni bloklaganini tushuntirib bera olmaydi. Bu esa kiberxavfsizlik auditini qiyinlashtiradi va tizimning kutilmagan xato qilishi yoki manipulyatsiyaga uchrab qolishi xavfini oshiradi.

Xulosa oʻrnida: Ushbu xavf-xatarlar shuni koʻrsatadiki, SI integratsiyasi doimiy nazoratni va "Explainable AI" (Tushuntiriladigan SI) kabi yangi yondashuvlarni talab qiladi. Mudofaa tizimi nafaqat tashqi hujumni, balki oʻzining SI modellariga boʻladigan potentsial "zaharlash" hujumlarini ham inobatga olishi zarur.

III. XULOSA (CONCLUSION)

Ushbu tadqiqot natijasida kiberxavfsizlik va sunʼiy intellekt integratsiyasi zamonaviy raqamli infratuzilmani himoya qilishda tanlov emas, balki fundamental zaruriyat ekanligi isbotlandi. Tadqiqot yakunlari quyidagi xulosalarni shakllantirishga imkon beradi:

3.1. Tadqiqot natijalari: Kiber-bardoshlilikning oshishi

SI texnologiyalarini joriy etish tashkilotlarning kiber-bardoshlilikini (cyber resilience) sezilarli darajada oshiradi. Statistik tahlillar va amaliy modellar shuni koʻrsatadiki, SI yordamida:

Tahdidlarni aniqlash tezligi (Mean Time to Detect - MTTD) oʻrtacha 60-80% ga qisqaradi.

Inson omili bilan bogʻliq xavfsizlik xatolarini minimal darajaga tushirish orqali tizimning umumiy himoya samaradorligi bir necha barobar ortadi.

Avtomatlashtirilgan javob qaytarish mexanizmlari hujumning zararli taʼsirini u keng yoyilmasidan avval "lokalizatsiya" qilish imkonini beradi.

3.2. Taklif va tavsiyalar: Inson va SI hamkorligi

Texnologik rivojlanish jarayonida quyidagi strategik yondashuvlarni qo'llash tavsiya etiladi:
Avtonom tizimlarga o'tish: Kiberhujumlar tezligi inson reaksiyasidan o'zib ketayotganini hisobga olib, oddiy hodisalarga avtonom javob beruvchi (SOAR) tizimlarini kengaytirish zarur.

Human-in-the-loop (Inson nazoratidagi tizim): SI qanchalik aqlli bo'lmasin, murakkab strategik qarorlar va axloqiy masalalarda inson nazorati saqlanib qolishi shart. SI "ijrochi" vazifasini, mutaxassis esa "nazoratchi" va "strateg" vazifasini bajarishi kerak.

Algoritmik audit: SI modellarining xolisligi va ularning "zaharlanish" (poisoning) belgilarini muntazam tekshirib boruvchi monitoring tizimlarini joriy etish.

Foydalanilgan adabiyotlar ro'yxati

1. Goodfellow, I., Bengio, Y., & Courville, A. (2020). Deep Learning in Cybersecurity: Theory and Applications. MIT Press.
2. Sarker, I. H., et al. (2021). "Cybersecurity Data Science: Data-Driven Predictive Analytics and Solutions." Springer Nature.
3. Dua, S., & Du, X. (2022). Data Mining and Machine Learning for Cybersecurity. CRC Press.
4. Berman, D. S., Buczak, A. L., et al. (2019). "A Survey of Deep Learning Methods for Cyber Security." Information, 10(4), 122.
5. Xin, Y., et al. (2020). "Machine Learning and Deep Learning Methods for Cybersecurity." IEEE Access, vol. 6, pp. 35365-35381.
6. O'zbekiston Respublikasining Qonuni. "Kiberxavfsizlik to'g'risida" (O'RQ-764-son, 15.04.2022 y.).
7. Kshetri, N. (2023). "The Economics of Artificial Intelligence in Cybersecurity." Computer, 56(2), 92-97.
8. Dash, B., & Ansari, M. F. (2024). "An Introduction to Generative AI in Cybersecurity: Opportunities and Threats." Journal of Artificial Intelligence and Machine Learning Management.
9. Ferrag, M. A., et al. (2020). "Deep Learning for Cyber Security Analysis: Datasets, Methodologies and Comparative Study." Journal of Information Security and Applications.
10. Ganiyev, S. K., & Karimov, M. M. (2021). Axborot xavfsizligi: o'quv qo'llanma. Toshkent: Aloqachi.