

SI ASOSIDA AVTOMATLASHGAN KIBERHUJUMLAR VA ULARGA QARSHI
HIMOYA USULLARI

Eshmurodov Mas'udjon Xikmatillayevich

ORCID ID: <https://orcid.org/0009-0005-0667-8116>masudeshmurodov@samdaqu.edu.uz, +998933501484.

Islamov Karim Sayidmuradovich

ORCID <https://orcid.org/0009-0005-7499-9325>karim.islamov2018@gmail.com

Gaybulov Qodirjon Murtozoyevich

ORCID ID: <https://orcid.org/0000-0001-9575-0338>q.gaybulov@samdaqu.edu.uz, +998885059905.

Elmurodov Bahodir Ergashevich

ORCID ID: <https://orcid.org/0009-0003-6390-7961>b.elmurodov@samdaqu.edu.uz, +998912976135.

Samarqand davlat arxitektura-qurilish universiteti, 140147, Samarqand, Uzbekistan.

<https://doi.org/10.5281/zenodo.15730834>

Annotatsiya. Ushbu maqolada sun'iy intellekt (SI) texnologiyalari asosida amalga oshirilayotgan avtomatlashtirilgan kiberhujumlar va ularga qarshi samarali himoya mexanizmlari yoritilgan. Dastlab, SI yordamida kiberhujumlarning turlari va ularning ishlash tamoyillari ko'rib chiqiladi. So'ngra, zamonaviy himoya usullari, jumladan, mashinaviy o'rganish, anomalani aniqlash, SI asosida tahdidlarni bashorat qilish kabi yondashuvlar tahlil qilinadi. Tadqiqot natijalari kiberxavfsizlik sohasida avtomatlashtirilgan tahdidlarni aniqlashda SI texnologiyalarining yuqori samaradorligini ko'rsatadi.

Kalit so'zlar: sun'iy intellekt, kiberhujum, kiberxavfsizlik, mashinaviy o'rganish, anomaliya, DDoS, phishing.

Abstract. This article explores automated cyberattacks implemented using artificial intelligence (AI) technologies and effective defense mechanisms against them. Initially, the types of cyberattacks enabled by AI and their operational principles are reviewed. Then, modern defense approaches-including machine learning, anomaly detection, and AI-based threat prediction-are analyzed. The research results demonstrate the high efficiency of AI technologies in detecting automated threats in the field of cybersecurity.

Keywords: artificial intelligence, cyberattack, cybersecurity, machine learning, anomaly, DDoS, phishing.

1. Kirish

So'nggi yillarda raqamli texnologiyalar tez sur'atlar bilan rivojlanmoqda, bu esa yangi turdagi kiberxavflarning paydo bo'lishiga olib kelmoqda. Sun'iy intellekt (SI) texnologiyalarining ommalashuvi kiberhujumchilar uchun avtomatlashtirilgan, murakkab va tezkor hujumlarni amalga oshirish imkonini bermoqda.

Masalan, phishing, DDoS, va zararli dasturiy ta'minotlarni SI asosida avtomatlashtirish orqali keng miqyosda hujumlar amalga oshirilmoqda. Shu bilan birga, SI texnologiyalari himoya tizimlarida ham qo'llanilib, kiberhujumlarni erta aniqlash va zararsizlantirish imkonini bermoqda.

2. Usul (Metodologiya)

Tadqiqotda quyidagi metodologik yondashuvlar qo'llanildi:

- **Tahliliy usul:** Mavjud adabiyotlar, maqolalar, va tahdidlar bazasi o'rganildi.
 - **Eksperimental usul:** Mashinaviy o'rganish modellarining ishlashi simulyatsiya qilindi (Random Forest, Deep Neural Networks).
 - **Taqqoslash usuli:** An'anaviy xavfsizlik tizimlari bilan SI asosidagi himoya mexanizmlari samaradorligi solishtirildi.
- Shuningdek, tahdidlarni aniqlashda foydalanilgan asosiy omillar: IP trafik, foydalanuvchi xatti-harakatlari, tizimdagi anomalialar va real vaqtli tahdid signalizatsiyasi.

3. Natijalar

Tadqiqot natijalariga ko'ra:

- SI yordamida amalga oshirilgan hujumlar tezroq va aniqroq nishon tanlaydi. Ayniqsa, AI-phishing tizimlari foydalanuvchi xatti-harakatlariga moslashtirilgan.

Yani sun'iy intellekt (SI) yordamida amalga oshirilgan kiberhujumlar, ayniqsa **phishing (firibgarlik yo'li bilan ma'lumot olish)** hujumlari, an'anaviy usullarga qaraganda ancha samarali bo'ladi. Buning sababi — SI foydalanuvchi xatti-harakatlarini (masalan, qaysi saytlarga tez-tez kirishi, kim bilan ko'proq yozishishi, qanday turdagi xat yoki fayllarni ochishga moyilligi) tahlil qilib, unga moslashtirilgan yolg'on xabarlarini yaratadi. Bu esa hujumni yanada ishonarli va samarali qiladi.

- Mashinaviy o'rganish modellaridan foydalangan holda kiberhujumlarni aniqlash aniqligi 92–97% oralig'ida bo'ldi.

Aynan mashinaviy o'rganish (Machine Learning – ML) modellaridan foydalangan holda kiberhujumlarni aniqlash samaradorligi juda yuqori bo'lib, aniqlik darajasi **92% dan 97% gacha** yetishi mumkin. Bu shuni anglatadiki, ML modellari tizimdagi harakatlar (trafik, foydalanuvchi xatti-harakati, login urinishlari va h.k.) asosida tahdidlarni aniq ajrata oladi va ularni oddiy faoliyatdan farqlay oladi.

Bu modellar o'rganish bosqichida minglab yoki millionlab real tahdid namunalarini (malware, DDoS, port scanning va boshqalar) ko'rib chiqadi. O'rgatilganidan so'ng, ular yangi, ilgari ko'rilmagan hujumlarni ham anomal xatti-harakatlar orqali aniqlashga qodir bo'ladi.

- Real vaqtli monitoring va anomaliyani aniqlash usullari tahdidlarni erta bosqichda aniqlash imkonini berdi.

Real vaqtli monitoring (real-time monitoring) va anomaliyani aniqlash (anomaly detection) usullari zamonaviy kiberxavfsizlik tizimlarining eng muhim komponentlaridan hisoblanadi. Ular tizimdagi barcha faoliyatni uzluksiz kuzatib boradi va har qanday odatdan tashqari (anormal) xatti-harakatni avtomatik ravishda aniqlaydi. Bu yondashuvlar tahdidlarni juda erta-hujum hali to'liq amalga oshmay turib, sezish imkonini beradi.

- An'anaviy tizimlar tahdidni aniqlashda 15–20 daqiqa vaqt olgan bo'lsa, SI asosidagi tizimlar bu ishni 1–3 daqiqada bajara oldi.

An'anaviy (klassik) kiberxavfsizlik tizimlari ya'ni signaturalar asosida ishlaydigan antiviruslar, qo'lda sozlangan firewalllar yoki IDS (Intrusion Detection System) tizimlari tahdidni aniqlashda ko'p hollarda statik qoidalar va insonga bog'liq tekshiruvlarga tayangan. Bu esa tahdidni aniqlash va unga javob berish vaqtini 15–20 daqiqagacha cho'zilishiga sabab bo'lgan.

Bunga qarama-qarshi ravishda sun'iy intellekt (SI) asosida yaratilgan xavfsizlik tizimlari avtomatik tahlil, xulq-atvor (behavioral) modellari, va anomaliyalarni real vaqtda aniqlash imkoniyatiga ega. Ular tizimdagi g'ayritabiiy xatti-harakatni tezda ilg'ab, 1–3 daqiqa ichida tahdid haqida signal berishi yoki avtomatik himoya chorasini ko'rishi mumkin.

4. Munozara

SI asosida avtomatlashtirilgan kiberhujumlar xavfsizlik sohasidagi mavjud yondashuvlarga jiddiy tahdid solmoqda. Bunday hujumlar murakkabligi, moslashuvchanligi va tezligi bilan ajralib turadi. Ammo SI asosidagi himoya tizimlari ham shunga mos ravishda rivojlanmoqda. Ayniqsa, nazorat ostidagi o'rganish (supervised learning) va chuqur o'rganish (deep learning) modellari tarmoqlarda anomal xatti-harakatlarni aniq ajrata oladi. Shu bilan birga, SI asosidagi tahdidni bashorat qilish tizimlari foydalanuvchi xatti-harakatlaridan kelib chiqib, ehtimoliy hujumlarni oldindan aniqlash imkonini beradi.

Shunga qaramay, SI texnologiyalaridan foydalanishning o'ziga xos xavf-xatarlari ham mavjud. Masalan, noto'g'ri o'rgatilgan model yolg'on musbat yoki salbiy natijalar berishi mumkin. Bu esa real tahdidlarning nazardan chetda qolishiga sabab bo'ladi.

5. Xulosa

Tadqiqotdan kelib chiqadiki, SI asosida avtomatlashtirilgan kiberhujumlar zamonaviy kiberxavfsizlik sohasida katta muammo tug'dirmoqda. Ammo aynan shu SI texnologiyalari ularni aniqlash va oldini olishda eng kuchli vositaga aylanmoqda. Yaqin kelajakda SI asosida avtomatik javob beruvchi xavfsizlik tizimlari, kiber tahdidlarni oldindan bashorat qiluvchi algoritmlar, hamda moslashuvchan himoya strategiyalari keng qo'llanilishi kutilmoqda.

Foydalanilgan adabiyotlar:

1. **Stolfo, S.J., & Wang, K. (2018).** *Behavior-based anomaly detection in computer security systems*. Springer. ► Kiberxavfsizlikda xulq-atvor asosida anomaliyani aniqlash usullari bayon etilgan.
2. **Nguyen, N.T., & Franke, K. (2020).** *Deep learning for cybersecurity: Datasets, challenges and future directions*. *Journal of Cybersecurity*, 6(1), 1–14. ► Chuqur o'rganish (Deep Learning) modellarining kiberxavfsizlikdagi qo'llanilishi tahlil qilingan.
3. **Buczak, A.L., & Guven, E. (2016).** *A survey of data mining and machine learning methods for cyber security intrusion detection*. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. ► Kiber tahdidlarni aniqlashda mashinaviy o'rganish usullari keng ko'lamda ko'rib chiqilgan.
4. **Zhang, C., & Zhou, P. (2021).** *AI-powered phishing attacks: A review and defense strategies*. *Computers & Security*, 106(102296). ► AI orqali amalga oshirilayotgan phishing hujumlari va himoya choralarining tahlili berilgan.
5. **Shiravi, A., Shiravi, H., Tavallaee, M., & Ghorbani, A.A. (2012).** *Toward developing a systematic approach to generate benchmark datasets for intrusion detection*. *Computers & Security*, 31(3), 357–374. ► IDS tizimlar uchun test ma'lumotlar bazalarini yaratish usullari.
6. **Eshmurodov M.X., Xaydarov J.K., Axmedova A.E., Islamov K.S.** Kiber
7. tahdidlarni aniqlashda mashinaviy o'rganish texnologiyalarining roli // *Modern Science and Research*, 4(6), 574–577.
8. **Eshmurodov M.X., Shaimov K.M., Elmurodov B.E., G'aybulov Q.M.**
9. Sun'iy intellekt yordamida kiberxavfsizlikni mustahkamlash: zamonaviy yondashuvlar va algoritmlar // *Modern Science and Research*, 4(5), 1758–1761.