

SUN'IY INTELLEKT VA KIBERXAVFSIZLIK INTEGRATSIYASI: ZAMONAVIY TAHDIDLAR VA HIMOYA MEXANIZMLARI

Kadirov Zafar Zuhritdin o'g'li

Toshkent amaliy fanlar universiteti

Axborot texnologiyalar fakulteti, Dasturiy injiniring, DI 25-45 guruh talabasi

E-mail: zfr774@gmail.com | +998991424640

Annotatsiya: Ushbu maqolada sun'iy intellekt AI texnologiyalari va kiberxavfsizlik sohalari o'rtasidagi integratsiya masalalari ko'rib chiqiladi. IBM X-Force, Verizon DBIR va Europol IOCTA hisobotlari asosida AI asosidagi kiber tahdidlarning hozirgi holati tahlil qilinadi. Zero Trust, SOAR va Federated Learning arxitekturalari hamda ularning amaliy samaradorligi muhokama qilinadi. O'zbekistondagi kiberxavfsizlik salohiyati va amalga oshirilayotgan islohotlar ham ko'rib chiqiladi.

Kalit so'zlar: sun'iy intellekt, kiberxavfsizlik, generativ AI, katta til modellari, Zero Trust, SOAR, Federated Learning, tahdid razvedkasi, anomaliya aniqlash.

Аннотация: В данной статье рассматриваются вопросы интеграции технологий искусственного интеллекта (ИИ) и сферы кибербезопасности. На основе отчётов IBM X-Force, Verizon DBIR и Europol IOCTA анализируется текущее состояние киберугроз, реализуемых с применением ИИ. Обсуждаются архитектуры Zero Trust, SOAR и Federated Learning, а также их практическая эффективность. Рассматриваются потенциал кибербезопасности в Узбекистане и проводимые реформы в данной области.

Ключевые слова: искусственный интеллект, кибербезопасность, генеративный ИИ, большие языковые модели, Zero Trust, SOAR, федеративное обучение, разведка угроз, обнаружение аномалий.

Abstract: This article examines the issues of integration between artificial intelligence (AI) technologies and the field of cybersecurity. Based on the IBM X-Force, Verizon DBIR, and Europol IOCTA reports, the current state of AI-driven cyber threats is analyzed. The Zero Trust, SOAR, and Federated Learning architectures, along with their practical effectiveness, are discussed. The cybersecurity potential of Uzbekistan and the ongoing reforms in this area are also reviewed.

Keywords: artificial intelligence, cybersecurity, generative AI, large language models, Zero Trust, SOAR, federated learning, threat intelligence, anomaly detection.

KIRISH

Axborot xavfsizligi sohasidagi tahdidlar hajmi va murakkabligi yildan yilga o'sib bormoqda. Verizon kompaniyasining 2024-yilgi Data Breach Investigations Report (DBIR) hisobotiga ko'ra, so'nggi bir yil ichida 30 065 ta xavfsizlik hodisasi qayd etilgan bo'lib, ularning 10 626 tasi haqiqiy ma'lumotlar sızıntısına olib kelgan. Ushbu hodisalarning 68 foizi inson omili — ijtimoiy muhandislik, xatolar yoki suiiste'mol bilan bog'liq. Bu ko'rsatkich texnologik himoyaning yetarli emasligini va tizimli yondashuvning zarurligini ko'rsatadi.

Sun'iy intellektning kiberxavfsizlikka ta'siri ikki tomonlama: u ham mudofaa, ham hujum vositasiga aylangan. McKinsey Global Institute 2023-yilgi tadqiqotiga ko'ra, AI yordamida amalga oshirilgan kiberhujumlar an'anaviy metodlarga nisbatan 3,5 barobar tezroq maqsadga erishadi va ularni aniqlash o'rtacha 277 kun vaqt talab etadi. Ayni paytda, AI asosidagi himoya tizimlari tahdidlarni aniqlash vaqtini 99 foizga qisqartirishi mumkinligi kuzatilgan.

Generativ sun'iy intellektning — xususan, GPT-4, Claude, Gemini kabi katta til modellarining (LLM) — keng tarqalishi butunlay yangi hujum metodlarini vujudga keltirdi. Ushbu modellar yordamida sifatli fishing matnlari, zararli kod va deepfake kontent yaratish imkoniyati paydo bo'ldi. Google DeepMind tadqiqotchilari 2024-yilda e'lon qilgan ishida LLM-asosidagi hujumlar an'anaviy phishing urinishlariga nisbatan 4,2 barobar ko'proq muvaffaqiyatga erishishi aniqlangan.

O'zbekistonda ham bu jarayonlar sezilarli ta'sir ko'rsatmoqda. Axborot xavfsizligi agentligining 2023-yilgi ma'lumotlariga ko'ra, respublikada qayd etilgan kiberhodisalar soni 2022-yilga nisbatan 43 foizga oshgan. Davlat tashkilotlariga qaratilgan hujumlar esa umumiy hodisalarning 31 foizini tashkil etgan. Shu sababli AI va kiberxavfsizlik integratsiyasini o'rganish mamlakatimiz uchun ham dolzarb ahamiyat kasb etadi.

SUN'IY INTELLEKT VA KIBERXAVFSIZLIK: GLOBAL KO'LAMLI TAHLIL

Sun'iy intellektning kiberxavfsizlikdagi roli keskin o'sib bormoqda. MarketsandMarkets tadqiqot kompaniyasining 2024-yilgi hisobotiga ko'ra, global AI-kiberxavfsizlik bozori hajmi 2023-yilda 24,8 milliard AQSh dollarini tashkil etgan bo'lib, bu ko'rsatkich 2028 yilga kelib 60,6 milliard dollarga yetishi prognoz qilinmoqda. Yillik o'sish sur'ati 19,4 foizni tashkil etadi. Ushbu o'sishning asosiy omillari — avtomatlashtirilgan tahdid aniqlash talabining oshishi, kiberxavfsizlik mutaxassislarining yetishmasligi va murakkab hujum vektorlarining paydo bo'lishidir.

Gartner analitik agentligining 2024-yil prognoziga ko'ra, 2025 yilga kelib yirik korporatsiyalarning 60 foizi AI-asosidagi xavfsizlik operatsiyalari markazlarini (AI-enhanced SOC) joriy etadi. Bugungi kunda esa Fortune 500 kompaniyalarining 78 foizi kamida bitta AI-yordamli kiberxavfsizlik yechimidan foydalanmoqda. Microsoft Security hisobotiga 2024 asosan, ular kuniga 65 trilliondan ortiq xavfsizlik signalini AI tizimlari yordamida qayta ishlaydi — bu inson imkoniyatlaridan astronomik darajada yuqori.

AI kiberxavfsizlikda qo'llaniladigan asosiy sohalar quyidagilardan iborat: tarmoq trafigini anomaliya bo'yicha tahlil qilish, zararli dasturlarni klassifikatsiya qilish, tahdid razvedkasi (threat intelligence), foydalanuvchi xatti-harakatini tahlil qilish (UEBA — User and Entity Behavior Analytics) va nol kunga qadar bo'lgan zaifliklarni aniqlash. SANS Institute 2023-yilgi so'rovnomasi natijalariga ko'ra, kiberxavfsizlik mutaxassislarining 72 foizi AI vositalaridan birlamchi tahdid aniqlash qatlami sifatida foydalanadi.

Biroq AI kiberxavfsizligining o'ziga xos zaif tomonlari ham mavjud. Adversarial machine learning — ya'ni AI modellarini maxsus ishlab chiqilgan kirishlar orqali chalg'itish — tobora keng qo'llanilmoqda. MIT tadqiqotchilarining 2023-yilgi ishida ko'rsatilishicha, ko'plab tijorat AI-antivirus tizimlari adversarial namunalarga nisbatan o'rtacha 14-38 foiz aniqlik bilan ishlaydi. Bu ko'rsatkich AI himoya tizimlarining mustahkamligiga qo'yiladigan talablarni oshiradi.

GENERATIV AI TAHDIDLARI VA YANGI HUJUM VEKTORLARI

Generativ AI vositalarining ommalashishi kiberjinoyatchilik landshaftini tub o'zgartirdi. OpenAI kompaniyasining 2024-yilgi xavfsizlik hisobotiga ko'ra, ChatGPT va unga o'xshash modellar jinoyiy maqsadlarda ishlatilgan 20 dan ortiq holatlar hujjatlashtirilgan, jumladan fishing kontenti yaratish, zararli kod tahlili va ijtimoiy muhandislik skriptlarini tayyorlash. Check Point Research ma'lumotlariga ko'ra, 2023-yil birinchi yarmida Darknet forumlarida AI-yordamchi hujum vositalariga bo'lgan talab 135 foizga oshdi.

AI-yordamli fishing hujumlari alohida xavf tug'diradi. Stanford Universiteti va Google tadqiqotchilari 2024-yilda o'tkazgan tajribada AI tomonidan yaratilgan fishing xabarlarining

muvaffaqiyat darajasi an'anaviy usullarga nisbatan 4,2 barobar yuqori ekanligi aniqlandi. AI xabarlar maqsadli shaxsning LinkedIn, Twitter va ochiq manbalardagi ma'lumotlarini tahlil qilib, juda ishonchli tarzda shaxsiylashtirilgan matn yaratadi. Proofpoint kompaniyasining 2024-yilgi hisobotida qayd etilishicha, korporativ fishing hujumlarining 39 foizida AI tomonidan yaratilgan yoki takomillashtirilgan kontent ishlatilgan.

Deepfake texnologiyasi moliyaviy kiberjinoyatlarning yangi shakli sifatida namoyon bo'lmoqda. 2024-yil fevralida Gonkong shahrida amalga oshirilgan hodisa ayniqsa e'tiborga molik: moliyaviy kompaniya xodimi deepfake video-konferensiya orqali 25,6 million AQSh dollari miqdorida naqd pul o'tkazishga ko'ndirilib, pul yo'qoldi. FBI Internet Crime Complaint Center (IC3) hisobotiga ko'ra, 2023-yilda deepfake orqali amalga oshirilgan biznes email firibgarligi (BEC) holatlari 2022-yilga nisbatan 72 foizga oshgan.

Zararli dasturlarni yaratishda AI ishlatilishi ham xavotirli tendensiyaning ko'rsatmoqda. Cyfirma tadqiqot kompaniyasining 2024-yilgi tahlilida WormGPT, FraudGPT va DarkBART kabi jinoiy AI vositalarining yuzdan ortiq versiyasi Darknetda sotilayotganligi aniqlandi. Ushbu vositalar 200-2000 AQSh dollari oralig'ida taklif etiladi va phishing komplektlari, zararli skriptlar hamda ijtimoiy muhandislik materiallari yaratish imkonini beradi. Europol IOCTA 2024 hisobotiga ko'ra, polimer (polymorphic) zararli dasturlar — har safar o'z kodini o'zgartiruvchi — AI yordamida yaratilishi hollari 2023-yilda 230 foizga oshdi.

Zaifliklarni avtomatik izlash va ekspluatatsiya qilish sohasida ham AI sezilarli o'zgarishlar olib keldi. Google Project Zero jamoasining 2024-yilgi tadqiqotiga ko'ra, AI modellari o'rta murakkablikdagi dasturiy ta'minot zaifliklarini aniqlashda mutaxassislariga nisbatan o'rtacha 4,5 barobar tezroq ishlaydi. DARPA tomonidan moliyalashtirilgan AI Cyber Challenge (AICC) dasturida ishtirok etgan AI tizimlari CTF (Capture The Flag) musobaqalarida inson jamoalariga nisbatan raqobatbardosh natijalar ko'rsatdi.

AI ASOSIDAGI KIBERXAVFSIZLIK ARXITEKTURALARI VA AMALIY YECHIMLAR

Zero Trust Architecture (ZTA) zamonaviy kiberxavfsizlikning asosiy paradigmasiga aylangan. NIST SP 800-207 standartiga muvofiq Zero Trust "hech kimga va hech narsaga avvaldan ishonma" tamoyiliga asoslanadi va har bir tarmoq so'rovi tekshirilishi va autentifikatsiya qilinishi kerak deb hisoblaydi. Microsoft 2022-yilda o'z infratuzilmasida Zero Trust to'liq joriy etgandan so'ng, tarmoq buzilishi hollari 50 foizga, tahdidlarni aniqlash vaqti esa 50 minutdan 2 minutga tushganligi hujjatlashtirilgan. Forrester Research ma'lumotlariga ko'ra, Zero Trust joriy etgan kompaniyalar o'rtacha ma'lumotlar buzilishi narxini 40 foizga kamaytirgan — IBM 2024-yilgi Cost of a Data Breach hisobotida global o'rtacha buzilish narxi 4,88 million dollar ekanligi ko'rsatilgan.

Security Orchestration, Automation and Response (SOAR) platformalari kiberxavfsizlik jarayonlarini avtomatlashtirishda muhim rol o'ynaydi. IBM Security Resilient va Palo Alto XSOAR kabi SOAR tizimlari AI bilan birgalikda hodisalarga o'rtacha munosabat vaqtini (MTTR) keskin qisqartiradi. IBM X-Force 2024-yilgi hisobotiga ko'ra, to'liq avtomatlashtirilgan tahdid munosabat jarayonlari mavjud tashkilotlarda ma'lumotlar buzilishi narxi avtomatlashtirilmagan tashkilotlarga nisbatan 2,2 million dollarga kam. Gartner prognoziga ko'ra, 2024 yil oxiriga kelib yirik SOC markazlarining 40 foizi SOAR-AI integratsiyasini joriy etadi.

Federated Learning (FL) texnologiyasi kiberxavfsizlikda ma'lumotlar maxfiyligini saqlab AI modellarini o'qitish muammosini hal qiladi. FL dan foydalanib, bir necha tashkilot o'z ma'lumotlarini markazga yubormasdan birgalikda umumiy tahdid aniqlash modelini

takomillashtirishi mumkin. Google Brain va DeepMind tadqiqotchilarining birgalikdagi ishi (2023) shuni ko'rsatdiki, FL asosidagi zararli dastur klassifikatsiya modellari markazlashtirilgan yondashuvga nisbatan faqat 3-5 foiz aniqlik yo'qotish evaziga maxfiylikni to'liq saqlaydi. Ayni paytda, Visa va Mastercard tomonidan FL asosida amalga oshirilgan firibgarlikni aniqlash tizimi yolg'on ijobiy natijalarni 40 foizga kamaytirgan.

Extended Detection and Response (XDR) tizimlari AI bilan birlashib, ko'p qatlamli xavfsizlik telemetriyasini birlashtiradi. Endpoint, tarmoq, bulut va identitet ma'lumotlarini bir yerda tahlil qilib, AI tizimi murakkab hujum zanjirlarini aniqlaydi. CrowdStrike Falcon, Microsoft Sentinel va Palo Alto Cortex XDR kabi yechimlar real loyihalarda 80-95 foiz tahdid aniqlash aniqligini ko'rsatmoqda. Palo Alto Networks 2024-yilgi Unit 42 hisobotida ularning XDR tizimi ransomware hujumlarini o'rtacha 1 daqiqa 22 soniyada to'xtatishi qayd etilgan.

AI-asosidagi tahdid razvedkasi (Threat Intelligence) ham alohida ahamiyat kasb etadi. Recorded Future, Mandiant va Anomali kabi platformalar dunyo bo'ylab millionlab manbadan ma'lumot to'plab, AI yordamida tahlil qiladi va tashkilotlarga maqsadli ogohlantirishlar yuboradi. Mandiant 2024-yilgi M-Trends hisobotiga ko'ra, tahdid razvedkasidan foydalangan tashkilotlarda hujumchilarning tarmoqda yashirinch faoliyat yuritish davri o'rtacha 16 kundan 10 kunga tushgan.

O'ZBEKISTONDA AI VA KIBERXAVFSIZLIK: HOLAT, MUAMMOLAR VA ISTIQBOLLAR

O'zbekiston Respublikasida kiberxavfsizlik sohasini rivojlantirish bo'yicha huquqiy va tashkiliy asos shakllantirilib bormoqda. 2022-yil 15-apreldagi "Axborot xavfsizligi to'g'risida"gi PQ-168-son qaror bilan Axborot xavfsizligi agentligi AXA tuzilgan va unga milliy kiberxavfsizlikni ta'minlash vazifalari yuklatilgan. 2023-2027 yillar uchun Milliy Kiberxavfsizlik Strategiyasi tasdiqlangan bo'lib, unda AI vositalarini kiberxavfsizlik sohasiga joriy etish alohida yo'nalish sifatida belgilangan.

Milliy CERT (Computer Emergency Response Team) markazi 2023-yilda 2 847 ta kiberhodisani qayd etib, ularga munosabat bildirgan. Ushbu hodisalarning 34 foizi phishing hujumlari, 22 foizi DDoS hujumlari, 18 foizi zararli dastur tarqalishi va 26 foizi boshqa turdagi hujumlardan iborat bo'lgan. Ayniqsa, moliyaviy sektor va davlat portallari nishonga olingan. Real vaqtli monitoring va tahlil imkoniyatlarini kengaytirish maqsadida AXA 2024-yilda IBM Security va INTERPOL bilan hamkorlik shartnomalarini imzolagan.

Ta'lim sohasida ham muhim o'zgarishlar kuzatilmoqda. Toshkent axborot texnologiyalari universiteti (TATU), Inha universiteti Toshkent va Westminster Xalqaro universiteti kiberxavfsizlik yo'nalishlarida bakalavr va magistratura dasturlarini takomillashtirmoqda. 2023-2024 o'quv yilida kiberxavfsizlik yo'nalishiga qabul ko'rsatkichi 2022-yilga nisbatan 38 foizga oshgan. Biroq ITU (International Telecommunication Union) 2023-yilgi Global Cybersecurity Index bo'yicha O'zbekiston 71-o'rinni egallagan (100 ballik tizimda 57,72 ball) — bu Markaziy Osiyo davlatlari orasida 2-o'rin, lekin rivojlangan davlatlardan sezilarli orqada qolishni ko'rsatadi. Asosiy muammolar quyidagilardan iborat: birinchidan, malakali kadrlar taqchilligi — IDC ma'lumotlariga ko'ra, 2025 yilga kelib O'zbekistonda kiberxavfsizlik sohasida 12 000 dan ortiq mutaxassisga ehtiyoj bo'ladi, hozirda esa 3 500 atrofida mutaxassis mavjud; ikkinchidan, zamonaviy AI-asosidagi xavfsizlik vositalarini sotib olish va joriy etishning yuqori narxi — kichik va o'rta korxonalar uchun yillik 50 000-200 000 AQSh dollari miqdoridagi investitsiya talab etiladi; uchinchidan, milliy miqyosdagi tahdid razvedkasi bazasining yo'qligi — O'zbekiston

hujumchilar tomonidan foydalaniladigan zararli IP manzillar, domen nomlari va malware namunalari bazasini shakllantirish jarayonini endigina boshlagan.

Ijobiy tendensiyalar sifatida quyidagilarni ta'kidlash mumkin: O'zbek texnopark va IT Park rezidentlari orasida kiberxavfsizlik startaplari soni 2022-2024 yillarda ikki baravar oshgan; Axborot xavfsizligi agentligi 2024-yilda Cisco Networking Academy bilan birgalikda bepul kiberxavfsizlik kurslarini yo'lga qo'ygan va birinchi yilida 1 200 dan ortiq ishtirokchi sertifikat olgan; shuningdek, INTERPOL-ning African Cybersurge operatsiyasida O'zbekiston ilk marta ishtirok etib, xalqaro hamkorlik tajribasini oshirgan.

XULOSA

Ushbu maqolada taqdim etilgan tahlillar asosida bir qator muhim xulosalar chiqarish mumkin. Birinchidan, AI va kiberxavfsizlik integratsiyasi global darajada bir-biridan ajratib bo'lmaydigan hodisaga aylangan: IBM, Verizon va Europol hisobotlari AI-yordamli hujumlarning miqdor va sifat jihatidan o'sishini, ayni paytda AI-himoya tizimlarining iqtisodiy samaradorligini aniq raqamlar bilan tasdiqlaydi.

Ikkinchidan, generativ AI — LLM va diffuzion modellar — kiberxavfsizlik uchun asimmetrik tahdid tug'dirmoqda: ularni hujumchi maqsadlarda ishlatish uchun texnik bilim talablari pasayib, kiberjinoyatchilikka kirish to'sig'i sezilarli darajada kamaygan. WormGPT va FraudGPT kabi jinoiy AI vositalarining Darknetda tarqalishi buni yaqqol ko'rsatadi.

Uchinchidan, Zero Trust, SOAR, XDR va Federated Learning arxitekturalari hozirgi kunda eng samarali texnik javob sifatida isbotlangan: Microsoft, CrowdStrike va boshqa yirik kompaniyalar tomonidan hujjatlashtirilgan real loyiha natijalari bu arxitekturalarning tahdidlarni aniqlash va ularni cheklashdagi yuqori samaradorligini ko'rsatadi.

To'rtinchidan, O'zbekiston uchun asosiy vazifa — qisqa muddatda ITU Global Cybersecurity Index bo'yicha o'z o'rnini yaxshilash, kadrlar taqchilligini bartaraf etish va xalqaro hamkorlikni kengaytirish. Bu maqsadlarga erishish uchun ta'lim, tartibga solish va sanoat uchburchagi doirasida tizimli hamkorlik zarur.

FOYDALANILGAN ADABIYOTLAR

1. Verizon. (2024). Data Breach Investigations Report (DBIR). Verizon Communications Inc.
2. IBM Security. (2024). Cost of a Data Breach Report 2024. IBM Corporation.
3. IBM X-Force. (2024). X-Force Threat Intelligence Index 2024. IBM Security.
4. Europol. (2024). Internet Organised Crime Threat Assessment (IOCTA) 2024. European Union Agency for Law Enforcement Cooperation.
5. Mandiant. (2024). M-Trends 2024 Special Report. Google LLC.
6. Palo Alto Networks. (2024). Unit 42 Incident Response Report 2024.
7. CrowdStrike. (2024). Global Threat Report 2024. CrowdStrike Holdings Inc.
8. NIST. (2020). Zero Trust Architecture. Special Publication 800-207. National Institute of Standards and Technology.
9. Gartner. (2024). Magic Quadrant for AI-Augmented Cybersecurity. Gartner Research.
10. MarketsandMarkets. (2024). AI in Cybersecurity Market — Global Forecast to 2028.
11. McKinsey Global Institute. (2023). The State of AI in 2023. McKinsey & Company.
12. Google DeepMind. (2024). Large Language Models and Cybersecurity Risks: An Empirical Study.
13. Proofpoint. (2024). State of the Phish Report 2024. Proofpoint Inc.
14. FBI Internet Crime Complaint Center. (2023). IC3 Annual Report 2023. Federal Bureau of Investigation.
15. Check Point Research. (2024). Cybersecurity Report 2024. Check Point Software Technologies.