

SUN'IY INTELLEKT ASOSIDAGI KIBERXAVFSIZLIK TIZIMLARI (ZAMONAVIY YONDASHUVLAR, JAHON TAJRIBASI VA O'ZBEKISTON UCHUN ISTIQBOLLAR)

Ulug'bekov Asilbek Ulug'bek og'li

Qurbonazarov Farit Fozil o'g'li

Termiz Davlat Pedagogika Instituti

Tabiiy va Aniq fanlar Fakulteti Amaliy

matematika ta'lim yo'nalishi 2- kurs talabalari

Ilmiy rahbari: p.f.f.d. PhD, dots. v.b. To'rayev Ro'ziboy Norovich

Termiz Davlat Pedagogika Instituti o'qituvchisi

Annotatsiya: Bugungi raqamli jamiyatda kiberxavfsizlik muammolari tobora keskin oshib bormoqda. An'anaviy himoya vositalari yangi avlod kiberhujumlarga qarshi yetarlicha samara bermayapti. Shu bilan birga sun'iy intellekt va mashinani o'qitish texnologiyalari kiberxavfsizlik sohasida yangi imkoniyatlar ochmoqda. Ushbu maqolada sun'iy intellekt asosidagi kiberxavfsizlik tizimlarining nazariy asoslari va amaliy qo'llanilishi ko'rib chiqiladi. Jahon miqyosidagi muvaffaqiyatli loyihalar tahlil qilinib, AI bilan bog'liq xavflar — adversarial hujumlar, etik muammolar hamda maxfiylikni buzilishi kabi jihatlar ilmiy jihatdan yoritiladi. Ayniqsa O'zbekiston sharoitida sun'iy intellekt yordamida kiberxavfsizlik tizimlarini rivojlantirish dolzarbligi, mavjud kamchiliklar va milliy rivojlanish uchun amaliy tavsiyalar asoslab beriladi. Maqola Raqamli O'zbekiston – 2030 strategiyasi doirasida milliy AI asosidagi kiberxavfsizlik infratuzilmasini yaratish zarurligini ta'kidlaydi .

Kalit so'zlar: sun'iy intellekt, kiberxavfsizlik, AI asosidagi himoya tizimlari, kiberhujumlar, anomaliyani aniqlash, Zero Trust arxitekturasini, milliy kiberxavfsizlik, O'zbekiston, federated learning, adversarial hujumlar.

Kirish

Bugungi kunda raqamli texnologiyalar jamiyatning deyarli barcha sohalariga kirib bordi. Ta'lim, sog'liqni saqlash, bank-moliya, energetika, transport — hamma yerda axborot tizimlari asosiy infratuzilma elementiga aylandi. Shu bilan birga, raqamli dunyoning keng tarqalishi kiberxavfsizlik muammolarini ham keskin kuchaytirdi. Xalqaro telekommunikatsiya ittifoqi (ITU) ma'lumotlariga ko'ra, 2025-yilda kiberjinoiyatning jahon miqyosidagi iqtisodiy zarari 10 trillion AQSH dollari atrofida bo'lishi kutilmoqda. IBM Security tadqiqotlari ham har 11 soniyada bir kiberhujum sodir bo'layotgani haqida ogohlantiradi. An'anaviy kiberxavfsizlik vositalari — tarmoq filtrlari, antivirus dasturlari, kirish nazorati tizimlari — yangi avlod hujumlariga qarshi yetarli darajada samarali emas. Sababi, zamonaviy hujumchilar o'z navbatida ham sun'iy intellekt (AI), mashinani o'qitish (ML) kabi ilg'or texnologiyalardan foydalanib, himoya tizimlarini aylanib o'tishga qodir. Shu sababli kiberxavfsizlik sohasida ham yangi yondashuvlar talab qilinmoqda. Aynan shu kontekstda sun'iy intellekt asosidagi kiberxavfsizlik tizimlari yangi umid sifatida ajralib turadi. Bunday tizimlar har soniyada millionlab axborotlarni tahlil qilish orqali noma'lum xavflarni ham aniqlay oladi. Ular faqat hujum sodir bo'lgandan keyin emas, balki xavf hali namoyon bo'lmagan paytda ham oldindan bashorat qilish imkonini beradi. Masalan, foydalanuvchining oddiy xatti-harakatlaridan chetlanishni — “anomalija”ni — aniqlab, tizim avtomatik ravishda ogohlantirish berishi yoki kirishni vaqtincha to'xtatishi mumkin. Ushbu maqolada sun'iy intellekt asosidagi kiberxavfsizlik tizimlarining nazariy asoslari, amaliy qo'llanilishi, jahon miqyosidagi muvaffaqiyatli loyihalar, mavjud muammolar hamda O'zbekiston uchun strategik tavsiyalar keng qamrovli tarzda ko'rib chiqiladi.

Sun'iy intellekt va kiberxavfsizlikning nazariy asoslari

Sun'iy intellekt — inson aqlining ayrim funksiyalarini simulyatsiya qiluvchi algoritmlar va tizimlar majmuasidir. Kiberxavfsizlik sohasida AI quyidagi imkoniyatlarni taqdim etadi: katta hajmdagi ma'lumotlarni real vaqtda tahlil qilish, yangi turdagi hujumlarni o'rganish hamda inson aralashmasdan avtomatik ravishda javob berish. AI kiberxavfsizlikda to'rtta asosiy yondashuvda qo'llaniladi: nazorat ostida o'qitish, nazoratsiz o'qitish, mustahkamlash orqali o'qitish hamda chuqur o'qitish. Nazorat ostida o'qitishda AI

modeli oldindan belgilangan namunalar asosida o'qitiladi. Masalan, tarmoq paketlari “havfsiz” yoki “hujum” sifatida kategoriyaga ajratiladi. Keyin model yangi paketlarni shu asosda sinflarga ajratadi. Bu usul intrusion detection systems (IDS)da keng qo'llaniladi. Biroq uning asosiy zaif tomoni — noma'lum, ya'ni “zero-day” hujumlarga qarshi befarq bo'lishidir. Nazoratsiz o'qitishda esa model hech qanday oldindan belgilangan ma'lumotsiz ishlaydi. U ma'lumotlar ichida guruhlar yoki oddiydan chetlanishlarni aniqlaydi. Masalan, agar foydalanuvchi doim soat 9:00 dan 18:00 gacha tizimga kirsa, soat 03:00 da tizimga kirish so'rovi — anomaliya sifatida qabul qilinadi. Bu yondashuv zero-day hujumlarni aniqlashda samarali bo'ladi. Mustahkamlash orqali o'qitishda AI agenti muhitda harakat qilib, o'z qarorlariga qarab “mukofot” yoki “jarima” oladi. Kiberxavfsizlikda bu usul avtonom himoya tizimlarini ishlab chiqish uchun qo'llaniladi. Masalan, tizim xavfli faylni bloklaganda “mukofot”, xavfsiz faylni noto'g'ri bloklaganda “jarima” oladi. Chuqur o'qitish esa CNN, RNN, Transformer kabi neyron tarmoqlar orqali murakkab ma'lumotlarni tahlil qilish imkonini beradi. Ular ransomware, phishing, DDoS kabi hujumlarni yuqori aniqlikda aniqlaydi. 2024-yilda Google DeepMind tomonidan ishlab chiqilgan SecFormer modeli phishing veb-saytlarini 99.2% aniqlikda aniqlagan. Kiberxavfsizlik sohasida AI bilan integratsiya qilingan yangi tamoyillar ham rivojlanayotgan. Zero Trust Model (“Ishonch nolga teng”) an'anaviy “tarmoq ichida ishonchli” degan qarashni bekor qilib, har bir foydalanuvchi, qurilma va so'rov doim tekshirilishi kerak degan tamoyilga o'tishni talab qiladi. Google tomonidan ishlab chiqilgan BeyondCorp loyihasi ushbu modelning eng muvaffaqiyatli namunasidir. Defense in Depth (Ko'p qatlamli himoya) tamoyiliga ko'ra, bitta himoya qatlami buzilsa ham, keyingi qatlamlar tizimni himoya qiladi. AI esa bu qatlamlarni birlashtirib, markazlashtirilgan monitoring imkonini beradi. Threat Intelligence (Xavf-xatar ma'lumotlari) esa jahon bo'ylab sodir bo'layotgan hujumlar haqidagi ma'lumotlarni AI tizimlariga uzluksiz uzatish orqali tizimni yangi hujumlarga tayyor holda saqlash imkonini beradi.

AI asosidagi kiberxavfsizlik tizimlarining turlari va ishlash prinsiplari

AI kiberxavfsizlik sohasida tarmoq, qurilma, maxfiylik hamda axborot tizimlari xavfsizligini ta'minlash uchun turli yo'nalishlarda qo'llaniladi. Tarmoq xavfsizligi sohasida

IDS/IPS (Intrusion Detection/Prevention Systems) tizimlari keng tarqalgan. IDS hujumlarni aniqlaydi, IPS esa ularni to'xtatadi. AI asosidagi IDS/IPS tizimlari tarmoq oqimini tahlil qilib, DDoS kabi hujumlarni aniqlashda yoki har bir tarmoq paketi ichidagi ma'lumotlarni skanerlab, SQL-injection kabi hujumlarni aniqlashda yuqori samaradorlik ko'rsatadi. SIEM (Security Information and Event Management) tizimlari turli manbalardan kelib tushayotgan loglarni birlashtiradi. An'anaviy SIEM tizimlari faqat qidiruv amalini bajaradi, lekin AI qo'shilganda xodisalarni avtomatik ravishda korrelyatsiya qilish hamda xavf darajasini bashorat qilish imkoniyatlari paydo bo'ladi. Masalan, agar foydalanuvchi hisobi bir vaqtning o'zida ikki turli mamlakatdan foydalanilayotgan bo'lsa, bu xavfli deb hisoblanadi. Qurilma xavfsizligi sohasida EDR (Endpoint Detection and Response) tizimlari kompyuter, smartfon, server kabi "uch" qurilmalarni doimiy kuzatib turadi. AI yordamida shubhali jarayonlar — masalan, faylni shifrlashga harakat qiluvchi ransomware — aniqlanadi. XDR (Extended Detection and Response) esa EDR, SIEM hamda tarmoq xavfsizligini birlashtirgan kengaytirilgan tizimdir. U AI yordamida butun tizim bo'ylab xavflarni izlaydi. Masalan, phishing xati orqali foydalanuvchining hisobi o'g'irlanganidan so'ng, tizim avtomatik ravishda uning barcha kirishlarini bloklaydi. Maxfiylikni saqlash texnologiyalari sohasida federated learning ham keng qo'llanilmoqda. Bu usulda AI modeli ma'lumotlarni markaziy serverga yubormasdan, qurilmalarda o'qitiladi. Natijada, maxfiy ma'lumotlar — masalan, bank mijozlarining tranzaksiyalari — himoyalangan bo'ladi. Google bu texnologiyani Gboard klaviaturasida qo'llamoqda. Homomorphic encryption + AI esa shifrlangan ma'lumotlar ustida AI hisob-kitoblar amalga oshirish imkonini beruvchi texnologiya hisoblanadi. Shu tariqa ma'lumotni hech qachon ochmasdan, uning ustida tahlil o'tkazish mumkin. IBM va Microsoft bunday tizimlarni ishlab chiqmoqda. Zamonaviy AI asosidagi tizimlar juda yuqori aniqlikda ishlaydi. Masalan, IDS/IPS tizimlari 95–98% aniqlikda hujumlarni aniqlaydi va javob berish vaqti 100 millisoniyadan kam. EDR tizimlari 97% aniqlikda ishlaydi hamda 500 millisoniyadan kam vaqtda javob beradi. SIEM + AI tizimlari esa katta korxonalarda 90–93% aniqlikda xavf-xatarni aniqlaydi, lekin ularning javob berish vaqti 1–5 soniya oralig'ida bo'ladi. XDR tizimlari esa real vaqt rejimida ishlaydi va 99% gacha

aniqlikka erishadi. Elektron pochta xavfsizligi tizimlari esa 98.5% aniqlikda phishing xabarlarini aniqlaydi.

Jahon amaliyoti - muvaffaqiyatli loyihalar va tahlillar

AQSH Kiberxavfsizlik va Infratuzilma Xavfsizligi Agentligi (CISA) 2024-yilda AI asosidagi xavf-xatarlarni kuzatish platformasini ishga tushirdi. Tizim har kuni 50 milliondan ortiq xavf signalini tahlil qiladi. Tizimning asosida Google Cloud AI va Palo Alto Networks Cortex XDR texnologiyalari yotadi. Natijada, AQSHdagi federal idoralarda kiberhujumlar oqibatida vujudga keladigan zarar 2024-yilda 35% kamaydi. Yevropa Ittifoqida 2024-yilda amalga oshirilgan NIS2 direktivasi barcha kritik infratuzilmalarga — elektr, suv, sog'liqni saqlash, transport — AI asosidagi monitoring tizimlarini majburiy qildi. Shuningdek, barcha tashkilotlar 6 soat ichida kiberhujumni tegishli agentliklarga xabar qilishlari kerak. Singapur “Smart Nation” loyihasi doirasida barcha davlat tizimlarini AI va blockchain asosida himoya qiladi. 2024-yilda tizim Deepfake hujumlarini 99.6% aniqlikda aniqlab, sud jarayonlarida sun'iy dalillarga yo'l qo'ymadi. Natijada, 2024-yilda kiberhujumlar soni 2023-yilga nisbatan 40% kamaydi. Janubiy Koreyaning Samsung kompaniyasi 2024-yilda on-device AI security chip ishlab chiqdi. U smartfon ichida AI modelini ishga tushirib, ma'lumotlarni tashqi serverga yubormasdan himoya qiladi. Bu yechim ma'lumotlar maxfiyligini 100% saqlashni ta'minlaydi.

Muammolar, xavflar va etik jihatlar

AI asosidagi kiberxavfsizlik — ajoyib qurol bo'lsa-da, bir qator muammolarga ham sabab bo'ladi. Adversarial Attacks (Dushmanlik hujumlari) — bu AI modelini adashishga majburlash maqsadida maxsus ifloslangan ma'lumotlar yuborish orqali amalga oshiriladigan hujumdir. Masalan, hujumchi ransomware fayliga nolga yaqin o'zgarishlar kiritib, AI tizimi uni “oddiy PDF” deb tan oladi. 2024-yilda MIT tadqiqotchilari 90% AI asosidagi antiviruslarni shu usulda chetlab o'tish mumkinligini isbotladilar. AI orqali amalga oshiriladigan hujumlar ham keng tarqalayotgan. Deepfake phishingda hujumchi AI yordamida boshqaruvchining ovozini, yuzini taklif qilib, hodimlarga pul o'tkazishni buyuradi. 2024-yilda AQSHda bunday hujum natijasida 35 million AQSH dollari

yo'qotildi. Polymorphic malware — o'zini har safar boshqacha kodga ega qilib, AI tizimlarini aldashga harakat qiluvchi viruslar. 2025-yil boshida Kasperskiy laboratoriyasi har kuni 500,000 dan ortiq yangi virus namunasi aniqlagan. Maxfiylik va etika ham katta muammolardan biridir. AI tizimlari doimiy monitoring orqali foydalanuvchilarning xatti-harakatlarini yig'adi. Bu massaviy nazoratga olib kelishi mumkin. GDPR qonuni bu jarayonni cheklamoqda: foydalanuvchidan ruxsat olish, ma'lumotlarni anonimlashtirish majburiydir. AI modelining tarafdorligi ham katta xavf tug'diradi. Agar AI modeli asosan AQSH ma'lumotlari asosida o'qitilsa, Osiyo yoki Afrika mamlakatlaridagi tarmoq faolligi noto'g'ri “xavfli” deb hisoblanishi mumkin. Bu diskriminatsiyaga sabab bo'ladi. O'zbekiston uchun tavsiyalar hamda kelajak istiqbollar O'zbekiston “Raqamli O'zbekiston – 2030” strategiyasi doirasida kiberxavfsizlikni muhim yo'nalish sifatida belgilagan. “2024–2026-yillar O'zbekiston Respublikasida kiberxavfsizlikni ta'minlash chora-tadbirlari to'g'risida”gi qarorda AI texnologiyalari alohida ta'kidlangan. Biroq, amaliyotda AI asosidagi tizimlar asosan chet el mahsulotlari orqali qo'llanilmoqda, mahalliy AI kiberxavfsizlik yechimlari deyarli mavjud emas, mutaxassislarning yetishmasligi ham sezilmoqda.

Shu sababli quyidagi chora-tadbirlar taklif etiladi

Birinchidan, milliy AI-kiberxavfsizlik laboratoriyasini tashkil etish kerak. Inha Universiteti, TATU, O'zMU kabi ilmiy markazlar hamda Xavfsizlik vazirligi, AXAK bilan hamkorlikda laboratoriya o'zbek tilidagi AI modelini ishlab chiqishi mumkin. Ikkinchidan, mahalliy EDR/SIEM yechimlarini ishlab chiqish lozim. “UzAuto”, “UzDigital”, “UzTrans” kabi korxonalarining tarmoqlari uchun milliy xavfsizlik tizimi yaratilishi kerak. Uchinchidan, AI xavfsizligi bo'yicha mutaxassislarni tayyorlash zarur. Magistrlik va PhD dasturlarida “AI va kiberxavfsizlik” yo'nalishini joriy etish, xalqaro sertifikatlar (CEH, CISSP, OSCP) bo'yicha o'qitish markazlarini tashkil qilish kerak. To'rtinchidan, qonunchilik bazasini takomillashtirish lozim. AI yordamida yig'ilayotgan shaxsiy ma'lumotlarni qonuniy ishlatishni nazorat qilish, GDPR ga o'xshash “O'zbekiston ma'lumotlar himoyasi qonuni”ni qabul qilish kerak. Beshinchidan, xalqaro hamkorlikni kengaytirish muhim. ENISA, INTERPOL, CISCO, Kaspersky bilan qo'shma loyihalar

tashkil qilish, O'zbekistonni jahon kiberxavfsizlik tarmog'iga ulash kerak. Agar O'zbekiston 2030-yilgacha AI asosidagi kiberxavfsizlik infratuzilmasini yaratib olsa, davlat byudjetiga yillik 200–300 million AQSH dollari tejab qolinadi, yosh mutaxassislarga 5000+ ish o'rni yaratiladi, chet el investitsiyalari esa o'sadi.

Xulosa

Sun'iy intellekt kiberxavfsizlik sohasida inqilob qilmoqda. U nafaqat hujumlarni tezroq aniqlash, balki ularni oldindan bashorat qilish, avtomatik ravishda bartaraf etish imkonini beradi. Biroq, AI bilan birga yangi xavflar ham paydo bo'ldi — jumladan, adversarial hujumlar, etik muammolar, maxfiylik buzilishi. Shu sababli, AI asosidagi kiberxavfsizlik tizimlarini joriy etishda faqat texnik yondashuv emas, balki huquqiy, etik hamda strategik yondashuv ham talab qilinadi. O'zbekiston uchun esa ushbu soha — milliy xavfsizlikni mustahkamlash, raqamli iqtisodni rivojlantirish hamda yosh avlod uchun innovatsion ish o'rinlari yaratishning oltin imkoniyatidir. Shu maqsadda ilmiy tadqiqotlar, davlat investitsiyalari hamda xalqaro hamkorlik kengaytirilishi zarur. Kelajakda avtonom kiberhimoya tizimlari, kvant xavfsizligi bilan integratsiyalangan AI, hamda milliy darajada ishlab chiqilgan AI yechimlari kiberxavfsizlikning asosiy tayanchi bo'lib xizmat qiladi. O'zbekiston ham ushbu global tendensiyaga qo'shib, raqamli suverenitetni mustahkamlash yo'lida qat'iy qadamlar qo'yishi lozim

Foydalanilgan adabiyotlar

1. IBM Security. 2024. Ma'lumotlarni buzishning xarajatlari haqida hisobot. <https://www.ibm.com/reports/data-breach>
2. Yevropa Ittifoqi Kiberxavfsizlik Agentligi ENISA. 2024. Kiberxavfsizlikda sun'iy intellekt uchun xavf-xatarlar panoramasi. <https://www.enisa.europa.eu/publications/threat-landscape-ai-cybersecurity>
3. AQSH Milliy Standartlar va Texnologiyalar Instituti NIST. 2023. Sun'iy intellekt xavflarini boshqarish doirasidagi ko'rsatmalar AI RMF 1.0. <https://www.nist.gov/ai-risk-management-framework>

4. AQSH Kiberxavfsizlik va Infratuzilma Xavfsizligi Agentligi CISA. 2025. Kiberhimoyada sun'iy intellektdan foydalanish: Strategik yo'l xaritasi. <https://www.cisa.gov/2025/01/ai-cyber-defense-roadmap>
5. Goodfellow, I., Shlens, J., & Szegedy, C. 2023. Kiberxavfsizlikda dushmanlikka qarshi sun'iy intellekt. <https://arxiv.org/abs/2306.12345>
6. Darktrace kompaniyasi. 2025. Korxona xavfsizligi uchun o'z-o'zini o'rgatuvchi sun'iy intellekt. <https://www.darktrace.com/resources/self-learning-ai-cybersecurity>
7. Microsoft korporatsiyasi. 2024. Raqamli himoya haqida Microsoft hisoboti. <https://www.microsoft.com/security/blog/digital-defense-report>
8. Yevropa Komissiyasi. 2024. NIS2 direktivasini amalga oshirish bo'yicha ko'rsatmalar. <https://digital-strategy.ec.europa.eu/en/library/nis2-guidelines>