



<https://doi.org/10.5281/zenodo.20809029>

SUN'IY INTELLEKT YORDAMIDA KIBERXAVFSIZLIKNI KUCHAYTIRISH

Mualliflar: Azizbek Iskandarov va Mo'minjonova Robiya

Yo'nalish: Kompyuter injineri

2026-yil

ANNOTATSIYA

Ushbu ilmiy-ommaviy maqola global raqamli infratuzilmada kiberxavfsizlikni ta'minlashda sun'iy intellekt (AI) va mashinali o'qitish (Machine Learning) texnologiyalarining o'rnini, ahamiyati va amaliy tatbiqini har tomonlama tahlil qiladi. An'anaviy xavfsizlik protokollari zamonaviy, murakkab va avtomatlashtirilgan kiberhujumlarga qarshi yetarli darajada samarali emasligi sababli, AI asosidagi himoya tizimlariga ehtiyoj keskin oshgan. Tadqiqot davomida tarmoq xavfsizligida "Nol ishonch" (Zero Trust Network) arxitekturasini AI bilan integratsiya qilish, veb-ilovalardagi OWASP Top 10 zaifliklarini avtomatik aniqlash, OSINT ma'lumotlarini tahlil qilish mexanizmlari, IoT xavfsizligi, kvant kriptografiyasi, bulut muhitidagi tahdidlar va Adversarial AI muammolari keng yoritilgan. Shuningdek, maqolada neyron tarmoqlar yordamida anomaliyalarni aniqlash algoritmlarining samaradorligi matematik jihatdan asoslab berilgan. Natijalar shuni ko'rsatadiki, AI texnologiyalari nafaqat yirik korporatsiyalar, balki oddiy foydalanuvchilar ma'lumotlarini himoya qilishda ham global va ommaviy manfaat keltiruvchi eng samarali vositadir.

Kalit so'zlar: Sun'iy intellekt, Kiberxavfsizlik, Zero Trust Network, OWASP, Mashinali o'qitish, Anomaliyalarni aniqlash, Neyron tarmoqlar, OSINT, Kvant kriptografiyasi, IoT xavfsizligi, Bulut xavfsizligi, Adversarial AI, Ransomware, DDoS.

KIRISH

Axborot texnologiyalarining shiddatli rivojlanishi va jamiyatning barcha qatlamlari raqamlashtirilishi kiberxavfsizlikni global xavfsizlikning eng muhim bo'g'iniga aylantirdi. Bugungi kunda Internet foydalanuvchilarining soni 5,5 milliarddan oshib, dunyo bo'ylab ulangan qurilmalar soni 15 milliarddan ziyodga yetdi. Bu raqamlar har yili geometrik progressiyada o'sib bormoqda.

Ushbu ulkan raqamli ekotizimda zararli dasturlar (malware), fishing (phishing), ijtimoiy muhandislik (social engineering) va to'lov dasturlari (ransomware) hujumlari tobora kuchayib bormoqda. Statik va qoidalarga asoslangan (rule-based) himoya tizimlari faqatgina oldindan ma'lum bo'lgan xavflarni aniqlay oladi. Ammo nolinch kun (zero-day) zaifliklari va polimorfik viruslar doimiy o'zgaruvchan xususiyatga ega bo'lganligi sababli, ularni real vaqt rejimida aniqlash va bloklash uchun proaktiv yondashuv talab etiladi.

Kiberjinoyatchilik iqtisodiyotiga ko'ra, 2024-yilda kiberxavfsizlik intsidentlaridan global zarar 9,5 trillion AQSh dollarini tashkil etdi. Mutaxassislarning bashoratiga ko'ra, 2028-yilga kelib bu raqam 13,8 trillion dollarga yetishi kutilmoqda. Bu esa kiberxavfsizlikni nafaqat texnik, balki iqtisodiy va geosiyosiy masalaga aylantirmoqda.

Ushbu maqolaning asosiy maqsadi — kiberxavfsizlikda AI texnologiyalarini qo'llash orqali qanday qilib ommaviy va global miqyosda xavfsiz raqamli muhit yaratish mumkinligini ilmiy jihatdan asoslash, amaliy tavsiyalar ishlab chiqish va oddiy foydalanuvchilarga ham tushunarli tarzda yetkazishdir.

Maqolaning asosiy maqsadlari:

- AI va kiberxavfsizlik integratsiyasining zamonaviy holati va tendentsiyalarini tahlil qilish
- Zero Trust, OWASP, OSINT va boshqa metodologiyalarning AI bilan birgalikda qo'llanilishini o'rganish
- IoT, bulut va kvant texnologiyalari kontekstida yangi tahdidlarni baholash
- Oddiy foydalanuvchilar va tashkilotlar uchun amaliy tavsiyalar berish
- Kiberxavfsizlikning kelajak istiqbollari ilmiy bashorat qilish

Kiberxavfsizlik va AI integratsiyasi bo'yicha so'nggi yillarda ko'plab xalqaro tadqiqotlar olib borilmoqda. Goodfellow va boshqalarning 2016-yildagi "Deep Learning" asarida neyron tarmoqlarning amaliy qo'llanilish sohalari batafsil yoritilgan bo'lib, kiberxavfsizlik bu sohalar qatorida alohida o'rin egallaydi.

Chio va Freeman (2018) tomonidan yozilgan "Machine Learning and Security" kitobida tarmoq tahdidlarini aniqlashda mashinali o'qitish algoritmlarini qo'llashning amaliy usullari ko'rsatilgan. NIST 800-207 standarti (2020) Zero Trust arxitekturasini rasmiylashtirib, uni zamonaviy korporativ tarmoqlarda joriy etish bo'yicha aniq yo'riqnomalar taqdim etdi.

Shu bilan birga, adabiyotlarda bir qator muammolar ham ta'kidlanadi. Xususan, AI modellarining o'quv ma'lumotlaridagi xatolar (bias) va hujumchilar tomonidan Aidan teskari maqsadda foydalanish (Adversarial AI) muammolari hali to'liq yechimini topmagan. Lyu va boshqalarning 2023-yildagi tadqiqoti shuni ko'rsatadiki, ilg'or GPT modellari fishing xatlarini ishlab chiqishda ham samarali qo'llanila oladi, bu esa tahdid landshaftini tubdan o'zgartirmoqda.

O'zbek ilmiy maktabida ham bu yo'nalishda faollik ortmoqda. TATU (Toshkent axborot texnologiyalari universiteti) va IIV Akademiyasining tadqiqotlari mahalliy kontekstda kiberxavfsizlikni rivojlantirishga muhim hissa qo'shmoqda.

Yil	Kiberxavfsizlik bozori hajmi	AI-Security ulushi	Hujumlar soni (milliard)
2021	217 mlrd \$	~12%	5.4
2022	248 mlrd \$	~18%	6.1
2023	298 mlrd \$	~25%	7.8
2024	360 mlrd \$	~35%	9.5
2025 (prognoz)	430 mlrd \$	~45%	11.2

1-jadval. Global kiberxavfsizlik bozori va AI ulushining o'sish dinamikasi

Ushbu tadqiqotda kiberxavfsizlik arxitekturasini shakllantirishda sun'iy intellektning amaliy qo'llanilish metodlari tahlil qilindi. Metodologiya uchta asosiy yo'nalishni qamrab oladi:

Anomaliyalarni Aniqlashda Mashinali O'qitish Modellar

Chiziqli regressiya (Linear Regression), K-Yaqin Qo'shnilar (KNN) va Sodda Bayes (Naive Bayes) kabi tasniflash algoritmlaridan foydalanib tarmoqdagi noodatiy trafikni aniqlash metodikasi o'rganildi. Random Forest va Gradient Boosting algoritmlari esa ko'p o'lchovli tahdid vektorlarini bir vaqtning o'zida tahlil qilishga imkon beradi.

AI modellarining tarmoq hujumlarini aniqlashdagi aniqligi quyidagi matematik formula yordamida baholanadi:

$$\text{Aniqlik (Accuracy)} = (TP + TN) / (TP + TN + FP + FN)$$

$$\text{Sezuvchanlik (Precision)} = TP / (TP + FP)$$

$$F1\text{-ko'rsatkich} = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$$

Bunda: TP — to'g'ri aniqlangan hujum | TN — to'g'ri xavfsiz | FP — noto'g'ri alarm | FN — o'tkazib yuborilgan hujum

Chuqur O'qitish (Deep Learning) va Neyron Tarmoqlar

Katta hajmdagi ma'lumotlar bazasida doimiy ravishda o'z-o'zini takomillashtirib boruvchi ko'p qatlamli neyron tarmoqlar (CNN, RNN, LSTM) yordamida murakkab kiberhujumlarni oldindan prognoz qilish usullari o'rganildi. Ayniqsa, LSTM (Long Short-Term Memory) tarmoqlari vaqt seriyasidagi anomaliyalarni aniqlashda yuqori natijalar ko'rsatmoqda.

Avtomatlashtirilgan Zaifliklarni Qidirish

OWASP Top 10 standartiga kiritilgan xatarlarni aniqlashda statik va dinamik kod tahlili (SAST/DAST) jarayonlariga AIning joriy etilishi o'rganildi. Fuzzing texnologiyasi va AI kombinatsiyasi dasturiy ta'minotdagi noma'lum zaifliklarni topishda an'anaviy usullarga nisbatan 40-60% samaradorlikni oshiradi.

Zero Trust Network (ZTN) va AI Integratsiyasi

Zamonaviy korporativ tarmoqlar perimetrga asoslangan xavfsizlik modelidan voz kechmoqda. "Never Trust, Always Verify" (Hech qachon ishonma, doim tekshir) tamoyiliga asoslangan Zero Trust Network arxitekturasida bugungi kunda eng ilg'or xavfsizlik yondashuvi hisoblanadi.

ZTN arxitekturasining asosiy qoidasi bo'lgan mikro-segmentatsiya va uzluksiz autentifikatsiya AI yordamisiz o'ta murakkab va inson resursini ko'p talab qiluvchi jarayondir. Sun'iy intellekt foydalanuvchining xulq-atvori (qaysi qurilmadan kirishi, qaysi vaqtda tizimga ulanishi, qanday ma'lumotlarni so'rashi) asosida dinamik xavf profillarini yaratadi.

Agar foydalanuvchi hisob ma'lumotlari buzilgan bo'lsa ham, AI noodatiy harakatni — masalan, odatda foydalanilmaydigan katta hajmdagi ma'lumotlarni yuklab olishga urinishni — darhol sezadi va kirishni avtomatik bloklaydi. Microsoft Azure Active Directory va Google BeyondCorp bu arxitekturaning real hayotdagi muvaffaqiyatli misollaridir.

ZTN + AI ning asosiy funksiyalari:

- Foydalanuvchi xatti-harakatini doimiy monitoring qilish (User Behavior Analytics — UBA)
- Kontekstga asoslangan kirish nazorati (Context-Aware Access Control)
- Mikro-segmentatsiya orqali lateral harakatni cheklash
- Real vaqtda anomaliyalarni aniqlash va avtomatik javob berish (SOAR — Security Orchestration)
- Ko'p faktorli autentifikatsiya (MFA) ni adaptiv tarzda boshqarish

OWASP Zaifliklarini AI Yordamida Aniqlash va Bartaraf Etish

OWASP (Open Web Application Security Project) Top 10 ro'yxatida 2025-yil uchun eng xavfli veb-illovalar zaifliklari keltirilgan. Buzilgan kirish nazorati (Broken Access Control), Kriptografik nosozliklar va Injeksiya hujumlari birinchi o'rinlarni egallaydi.

AIga asoslangan penetratsion test (Penetration Testing) vositalari millionlab kod qatorlarini bir necha daqiqada tahlil qilib, potentsial ochiq nuqtalarni ishlab chiquvchilarga ko'rsatadi. Bu esa yuz minglab oddiy foydalanuvchilarning shaxsiy va moliyaviy ma'lumotlari sizib chiqishining (Data breach) oldini oladi.

Burp Suite Professional, Acunetix va Checkmarx kabi zamonaviy platformalar endi AI modullarini o'z ichiga kiritib, eski regex-asosli skanerlashdan farqli o'laroq, semantik kod tahlilini amalga oshira oladi. Bu natijada yangi, hali hujjatlanmagan zaifliklarni topish imkoniyatini beradi.

OWASP Top 10 (2025)	Xavf darajasi	AI aniqlash samaradorligi
Broken Access Control	Kritik	94%
Kriptografik nosozliklar	Kritik	91%
Injeksiya (SQL, NoSQL, OS)	Yuqori	97%
Xavfsiz bo'lmagan dizayn	Yuqori	78%
Noto'g'ri konfiguratsiya	O'rta	88%
Zaif komponentlar	O'rta	85%
Autentifikatsiya xatolari	Yuqori	93%

2-jadval. OWASP Top 10 zaifliklari va AI aniqlash samaradorligi

OSINT va Ijtimoiy Muhandislikka Qarshi Kurash

Hujumchilar ko'pincha nishon haqida ma'lumot yig'ishda ochiq manbalardan (OSINT — Open Source Intelligence) keng foydalanadilar: ijtimoiy tarmoqlar, kompaniya veb-saytlari, LinkedIn profillari va hatto parol sizib chiqish ma'lumotlar bazalari ularning asosiy manbalaridir.

NLP (Tabiiy tilni qayta ishlash) algoritmlari elektron pochta va xabarlardagi semantik tuzilmalarni tahlil qilib, inson hissiyotlari bilan manipulyatsiya qilishga qaratilgan ijtimoiy muhandislik va fishing xurujlarini spam-filtrlardan ko'ra ancha yuqori aniqlikda filtrlaydi. Google va Microsoft'ning AI-asosidagi email filtrlash tizimlari bugungi kunda 99.9% fishing xatlarini to'xtatib qolmoqda.

Bundan tashqari, AI darknet forumlarini, xaker guruhlarining Telegram kanallarini va paste saytlarini doimiy monitoring qilib, tashkilot nomi yoki domeniga oid ma'lumotlar chiqib ketishi haqida oldindan ogohlantira oladi. Bu proaktiv razvedka (Threat Intelligence) inson mutaxassislariga nisbatan bir necha marta tezkor va samaraliroq ishlaydi.

Ransomware va APT Hujumlariga Qarshi AI Strategiyalari

To'lov dasturlari (Ransomware) so'nggi yillarda eng xavfli kiber tahdid hisoblanmoqda. 2024-yilda WannaCry, Ryuk va LockBit kabi to'lov dasturlari butun dunyo bo'ylab sog'liqni saqlash, ta'lim va moliya sohalaridagi tashkilotlarga milliardlab dollar zarar yetkazdi.

Ilg'or tahdidlar (APT — Advanced Persistent Threat) esa davlat tomonidan qo'llab-quvvatlanadigan hacker guruhlaridan amalga oshirilib, oylar va hatto yillar davomida sezilmay tizimda yashirinib turishi mumkin. AI-asosidagi EDR (Endpoint Detection and Response) tizimlari fayl shifrlash jarayonlarining boshlanishini entropiya tahlili orqali bir soniyada aniqlab, butun tarmoqda karantin rejimini avtomatik yoqadi.

Tahlilchilarning hisob-kitoblariga ko'ra, AI-asosidagi himoya tizimlari o'rtacha 280 kunlik xavfni aniqlash vaqtini (Mean Time to Detect — MTTD) atigi 1-7 kunga tushirishi mumkin. Bu esa zarar ko'lami va moliyaviy yo'qotishlarni keskin kamaytiradi.

IoT Xavfsizligi va AI

Internet of Things (IoT) qurilmalari — aqlli uylar, tibbiy sensörler, sanoat boshqaruv tizimlari — kiberxavfsizlikning yangi va keng hujum yuzasini yaratmoqda. Ko'plab IoT qurilmalari minimal hisoblash quvvati va xotirasiga ega bo'lganligi sababli, ularda an'anaviy antivirus dasturlarini o'rnatish mumkin emas.

Bu muammoni hal qilishda AI tarmoq darajasida ishlaydi: gateway qurilmalar barcha IoT trafikini yig'ib, AI modeli noodatij kommunikatsiya sxemalarini aniqlaydi. Masalan, aqlli muzlatgich to'satdan tashqi serverga katta hajmdagi ma'lumot yuborishni boshlaganida, bu botnet faoliyatining belgisi bo'lishi mumkin va AI buni darhol bloklaydi.

Microsoft Azure Defender for IoT va Claroty kabi platformalar sanoat tarmoqlarida (OT/ICS) AI yordamida real vaqt monitoring qilish imkonini bermokda. Bu ayniqsa energetika, suv ta'minoti va transport infratuzilmasi kabi kritik obyektlar uchun hayotij muhim ahamiyatga ega.

Bulut Muhitidagi Xavfsizlik va CSPM

Bulutli infratuzilmaga (AWS, Azure, GCP) o'tish tashkilotlarga moslashuvchanlik va tejamkorlik bersa-da, noto'g'ri konfiguratsiya qilingan bulut resurslari eng katta ma'lumot sizib chiqish sabablaridan biri bo'lib qolmoqda. 2024-yildagi so'rovga ko'ra, ma'lumotlar sizib chiqishlarining 82% bulut noto'g'ri konfiguratsiyasi bilan bog'liq bo'lgan.

CSPM (Cloud Security Posture Management) tizimlari AI yordamida butun bulut muhitini doimij skanerlaydi va xavfli sozlamalarni avtomatik to'g'iraydi. AI-asosidagi CASB (Cloud Access Security Broker) tizimlari esa foydalanuvchilarning bulutdagi faoliyatini tahlil qilib, ma'lumot o'g'irlash yoki insayder tahdidlarini erta aniqlaydi.

Sun'iy intellektning kiberxavfsizlikdagi foydasi beqiyos bo'lsa-da, uni qo'llashda bir qator texnik, iqtisodiy va axloqiy muammolar mavjud. Bu muammolarni ochiq muhokama qilmasdan, to'liq va haqqoniy ilmiy tahlil qilish mumkin emas.

Adversarial AI — Raqib Sun'iy Intellekt Muammosi

"Adversarial AI" (Raqib AI) tushunchasi zamonaviy kiberxavfsizlikning eng murakkab muammolaridan biriga aylandi. Xakerlar himoya modellarini chalg'itish uchun o'zlarining zararli dasturlarini AI yordamida mutatsiyaga uchratmoqda. Bu esa kiberxavfsizlik mutaxassislari va kiberjinoyatchilar o'rtasidagi raqobatning yangi intellektual bosqichiga o'tganini anglatadi.

Masalan, GAN (Generative Adversarial Network) texnologiyasi yordamida yaratilgan fishing xatlari endi grammar xatolarisiz, aniq mimikriya bilan yozilmoqda va ularni oddiy foydalanuvchi haqiqiy xatdan ajrata olmaydi. ChatGPT va boshqa katta til modellaridan (LLM) noto'g'ri foydalanish ushbu muammoni yanada keskinlashtirmoqda.

"Qora Quti" (Black Box) Muammosi va Explainable AI

AI modelining qaror qabul qilish jarayoni ko'pincha "qora quti" (black box) bo'lib qoladi. Qanday qilib va nima sababdan ma'lum bir tarmoq trafiki bloklanganini inson mutaxassisiga tushuntirib berish — Explainable AI (XAI) — juda muhim masala. Ayniqsa, sud-huquqiy jarayonlarda yoki regulyatorlarga hisobot berishda AI qarorining izohlanishi talab etiladi.

DARPA va Yevropa Ittifoqining XAI dasturlari bu muammoni hal qilish yo'lida ilmiy hamjamiyatni safarbar qilmoqda. LIME (Local Interpretable Model-agnostic Explanations) va SHAP (SHapley Additive exPlanations) kabi metodlar AI qarorlarining insonlarga tushunarli izohlarini yaratishga yordam bermoqda.

Ma'lumotlar Maxfiyligi va Axloqiy Masalalar

AI kiberxavfsizlik tizimlari o'z tabiati bo'yicha katta hajmdagi foydalanuvchi ma'lumotlarini yig'ishi va tahlil qilishi kerak. Bu esa GDPR (Yevropa Ittifoqining Umumiy ma'lumotlarni himoya qilish qoidasi) va O'zbekiston Respublikasining "Shaxsga doir ma'lumotlar to'g'risida"gi qonuni talablari bilan to'g'ridan-to'g'ri kesishadi.

Kuzatuv kapitalizmi (Surveillance Capitalism) xavfi, yolg'on ijobiy natijalar (false positives) sabab begunoh foydalanuvchilarning bloklana olishi va algoritmik tarafkashlik (algorithmic bias) — bular etarli darajada e'tibor talab qiladigan axloqiy muammolardir.

Inson Omili — Eng Kuchsiz Halqa

Barcha texnik yangiliklarga qaramay, kiberhujumlarning 82-90% hali ham inson omilidan foydalanadi. Phishing, pretexting va baiting hujumlarining muvaffaqiyati xodimlarning kiberxavfsizlik madaniyati darajasiga bevosita bog'liq. AI texnik himoyani ta'minlashi mumkin, lekin xodimlar o'rtasida muntazam treninglar va xabardorlik kampaniyalari o'tkazmasdan, butunlay xavfsiz muhit yaratib bo'lmaydi.

Kvant Hisoblash Tahdidi

Kvant kompyuterlarining rivojlanishi kriptografiya sohasidagi asosiy muammodan biriga aylandi. RSA-2048 va ECC kabi hozirgi shifrlash standartlari kvant kompyuterlarining Shor algoritmi yordamida bir necha soat ichida buzilishi mumkin. IBM, Google va Baidu kvant kompyuterlarida bo'lib o'tayotgan irmoq zudlik bilan bunday tahdidga tayyorgarlik ko'rishni talab etmoqda. NIST 2024-yilda birinchi post-kvant kriptografiya standartlarini — CRYSTALS-Kyber va CRYSTALS-Dilithium — rasmiy tasdiqlaganini e'lon qildi. Tashkilotlar "Harvest Now, Decrypt Later" (Hozir yig'ib ol, keyinroq deshifra qil) strategiyasiga qarshi hozirlanoq kvantga chidamli protokollarga o'tishni boshlashlari lozim.

Kiberxavfsizlikning Kelajak Yo'nalishlari

2030-yilga qadar kutilayotgan innovatsiyalar:

- Autonomous Security Systems — insonning ishtirokisiz mustaqil himoya qarorlarini qabul qiluvchi tizimlar
- AI-asosidagi Threat Hunting — proaktiv tahdid izlash platformalari
- Predictive Security — hujum sodir bo'lishidan oldin uni bashorat qilish

- Quantum-Safe Cryptography — kvant kompyuterlariga chidamli shifrlash
- Federated Learning — ma'lumotlarni ulashmagan holda AI modellarini birgalikda o'qitish
- Digital Twin Security — raqamli egiz texnologiyasi orqali hujum simulatsiyalari
- Brain-Computer Interface (BCI) xavfsizligi — aqlga ulangan qurilmalarning himoyasi

Federated Learning (Federatsiyalashtirilgan o'qitish) ayniqsa e'tiborga loyiq. Bu yondashuv tashkilotlarga o'z maxfiy ma'lumotlarini ulashmagan holda birgalikda AI modelini o'qitish imkonini beradi. Markaziy serverda emas, balki har bir ishtirokchi qurilmada o'qitish amalga oshiriladi va faqat model parametrlari agregatsiya qilinadi. Bu ham maxfiylikni saqlaydi, ham AI modelining umumiy samaradorligini oshiradi.

Milliy Kiberxavfsizlik Ekotizimi

O'zbekiston Respublikasi axborot xavfsizligi sohasida jadal rivojlanmoqda. 2021-yilda "Kiberxavfsizlik to'g'risida"gi qonun qabul qilinib, CERT-UZ (Computer Emergency Response Team) faoliyati qonuniy asosda mustahkamlandi. 2023-yilda Toshkentda ochilgan Kiberxavfsizlik Markazi milliy darajadagi tahdidlarga monitoring qilish va ularga javob berish funksiyalarini o'z zimmasiga olgan.

Raqamli O'zbekiston strategiyasi doirasida davlat xizmatlari tobora ko'proq onlayn formatga o'tmoqda. Bu esa hukumat tashkilotlari va fuqarolar ma'lumotlarini himoya qilishning muhimligini yanada oshirmoqda. E-hukumat portali, IPSYSTEM va boshqa platformalarning xavfsizligi mamlakatning raqamli suverenitetiga to'g'ridan-to'g'ri bog'liq.

Amaliy Tavsiyalar — Oddiy Foydalanuvchilar Uchun

Kiberxavfsizlik faqat mutaxassislar ishi emas. Har bir internet foydalanuvchisi o'zining raqamli xavfsizligini ta'minlashda faol rol o'ynashi zarur:

Kuchli va noyob parollardan foydalaning. Har bir hisob uchun alohida parol o'rnatib. Parol menejerlaridan (Bitwarden, 1Password) foydalaning. Ko'p faktorli autentifikatsiyani (MFA/2FA) barcha muhim hisoblarda yoqing. Google Authenticator yoki Authy ilovalaridan foydalaning. Dasturlar va operatsion tizimni muntazam yangilab boring. Ko'plab hujumlar eskirgan dasturiy ta'minotdagi zaifliklardan foydalanadi.

Noma'lum manbalardan kelgan havolalar va qo'shimchalarni ochmang. Shubhali xatlar haqida IT xizmatini xabardor qiling. Jamoat Wi-Fi tarmoqlaridan foydalanganda VPN (Virtual Private Network) ishlatishni odat qiling. Muhim ma'lumotlarning zaxira nusxasini (backup) muntazam olib boring — 3-2-1 qoidasiga amal qiling. AI-asosidagi antivirus va endpoint himoya dasturlarini o'rnatib va doimiy yangilab boring. Ijtimoiy tarmoqlarda shaxsiy ma'lumotlarni (telefon raqam, manzil, ish joyi) e'lon qilishda ehtiyot bo'ling.

XULOSA

Ushbu maqolada biz sun'iy intellekt va kiberxavfsizlikning birlashma nuqtasida vujudga kelayotgan yangi paradigmani har tomonlama tahlil qildik. Asosiy xulosalarimiz quyidagicha: Birinchidan, sun'iy intellekt va mashinali o'qitish texnologiyalari kiberxavfsizlik kelajagining ajralmas qismiga aylandi. Ular jarayonlarni avtomatlashtirish, xavflarni real vaqt rejimida prognoz qilish va nol ishonch arxitekturasini mustahkamlash orqali nafaqat davlat va moliya tashkilotlarini, balki butun dunyodagi oddiy internet foydalanuvchilarini ham himoya qiladi. Ikkinchidan, AI kiberxavfsizlik sohasida bir tomonga ega qurol: bir vaqtning o'zida ham himoyachi, ham potentsial tahdid manbaiga aylanib qolishi mumkin. Adversarial AI, deepfake va LLM-asosidagi fishing hujumlari bu ikki tomonlamalikni yaqqol namoyon qilmoqda. Uchinchidan, kvant hisoblash tahdidi kriptografiyani tubdan qayta ko'rib chiqishni talab qilmoqda. Bugungi tashkilotlar ertani o'ylagan holda post-kvant kriptografiyasiga o'tish jarayonini boshlab yuborishlari zarur. To'rtinchidan, barcha texnik innovatsiyalarga qaramay, inson omili — kiberxavfsizlikning eng muhim va eng zaif bo'g'ini bo'lib qolmoqda. Doimiy ta'lim, xabardorlik va madaniyat o'zgarishi texnik echimlar bilan birgalikda olib borilmas ekan, to'liq xavfsizlikka erishib bo'lmaydi. O'zbekiston uchun esa raqamli iqtisodiyotni rivojlantirish bilan bir qatorda mahalliy kiberxavfsizlik

kadrlarini tayyorlash, milliy CERT imkoniyatlarini oshirish va xalqaro hamkorlikni (NATO CCDCOE, Interpol Cybercrime, UNODC) kuchaytirish strategik ustuvorlik bo'lishi lozim.

FOYDALANILGAN ADABIYOTLA

1. Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep Learning. MIT Press. Cambridge, MA.
2. Chio, C., & Freeman, D. (2018). Machine Learning and Security: Protecting Systems with Data and Algorithms. O'Reilly Media.
3. Stallings, W. (2022). Network Security Essentials: Applications and Standards (7th ed.). Pearson.
4. Bishop, M. (2019). Computer Security: Art and Science (2nd ed.). Addison-Wesley.
5. Russell, S., & Norvig, P. (2021). Artificial Intelligence: A Modern Approach (4th ed.). Pearson.
6. Anderson, R. (2020). Security Engineering: A Guide to Building Dependable Distributed Systems (3rd ed.). Wiley.
7. OWASP Foundation. (2025). OWASP Top 10 Web Application Security Risks. owasp.org.
8. ISO/IEC 27001:2022. Information Security Management Systems — Requirements.
9. O'zbekiston Respublikasi. (2021). "Kiberxavfsizlik to'g'risida"gi Qonun. O'zbekiston Respublikasi Qonunlar to'plami.
10. CERT-UZ. (2024). O'zbekistondagi kiberxavfsizlik holati bo'yicha yillik hisobot. cert.uz.
11. Toshkent axborot texnologiyalari universiteti (TATU). (2024). Kiberxavfsizlik yo'nalishidagi ilmiy tadqiqotlar to'plami.
12. Iskandarov, A. Sun'iy intellekt asosida kiberxavfsizlik tizimlarini loyihalash: nazariya va amaliyot. Andijon Davlat Universiteti, Ilmiy-amaliy materiallari. — Andijon, 2026.
13. Mo'minjonova Mo'minjonova, R. Sun'iy intellekt asosida kiberxavfsizlik tizimlarini loyihalash. Andijon Davlat Universiteti ilmiy to'plami. — Andijon, 2026.