

<https://doi.org/10.36719/2663-4619/112/146-153>

Sevil Haqverdiyeva

Gəncə, Azərbaycan

<https://orcid.org/0009-0000-1238-8338>

sevil.haqverdiyeva@gmail.com

Ceyhun Həsənov

Bakı, Azərbaycan

<https://orcid.org/0009-0007-4269-7327>

chasanoff@gmail.com

Elnarə Kazımova

Gəncə, Azərbaycan

<https://orcid.org/0009-0006-8994-5351>

kazimovaelnara617@gmail.com

Nailə Həsənova

Gəncə, Azərbaycan

<https://orcid.org/0009-0003-1549-1690>

nailahesenova1974@gmail.com

Anar Quliyev

Gəncə, Azərbaycan

<https://orcid.org/0009-0002-5549-4063>

anar_quliyev74@mail.ru

Kibertəhlükəsizlikdə süni intellekt və maşın öyrənməsi

Xülasə

Süni intellekt (Sİ) və onun alt sahəsi olan maşın öyrənməsi (MÖ), kompüterlərin insan kimi düşünmə, öyrənmə və qərar qəbul etmə bacarığını təmin edən texnologiyalardır. MÖ-nün əsas alqoritmləri nəzarətli öyrənmə, nəzarətsiz öyrənmə və dərin öyrənmədir. Nəzarətli öyrənmə, etiketlenmiş verilənlər üzərində işləyərək modellərin təlimini təmin edir. Nəzarətsiz öyrənmə isə etiketlenməmiş verilənlərdə gizli strukturları aşkarlayır. Dərin öyrənmə isə insan beynini təqlid edən çoxqatlı sinir şəbəkələrinə əsaslanır və böyük verilənlərlə işləyərək mürəkkəb problemləri həll edir.

Kibertəhlükəsizlik sahəsində Sİ və MÖ alqoritmləri təhlükələrin aşkarlanması, anomal nümunələrin müəyyən edilməsi və avtomatlaşdırılmış cavab sistemləri kimi mühüm funksiyaları yerinə yetirir. Məsələn, dərin öyrənmə alqoritmləri şəbəkə trafikində anomaliyaları aşkar edərək DDoS hücumlarının qarşısını ala bilər. Bununla yanaşı, Sİ sistemləri yalançı pozitivləri azaltmaq, təhlükələrə proqnozlaşdırmaq və təhlükəsizlik tədbirlərini avtomatlaşdırmaq imkanı yaradır.

Lakin, Sİ-nin kibertəhlükəsizlikdə tətbiqi zamanı insan amili hələ də əvəzolunmazdır. Strateji düşüncə, etik qərarlar və yaradıcı problem həlli kimi sahələrdə insan mütəxəssislərin rolu böyükdür. Gələcəkdə kibertəhlükəsizlik sahəsində insan və Sİ arasında əməkdaşlıq daha da artacaq, lakin məlumat keyfiyyəti, bacarıq boşluğu və etik məsələlər kimi çətinliklər də qarşıya çıxacaq. Beləliklə, kibertəhlükəsizliyin gələcəyi insan və Sİ-nin sinerjisindən asılı olacaq.

Bu məqalədə Sİ və MÖ-nün kibertəhlükəsizliyin müxtəlif sahələrində tətbiqi, üstünlükləri və çətinlikləri araşdırılacaq.

Açar sözlər: informasiya təhlükəsizliyi, kibertəhlükəsizlik, kiberhücumlar, süni intellekt, maşın öyrənməsi

Sevil Haqverdiyeva

Ganja, Azerbaijan

<https://orcid.org/0009-0000-1238-8338>

sevil.haqverdiyeva@gmail.com

Jeyhun Hasanov

Baku, Azerbaijan

<https://orcid.org/0009-0007-4269-7327>

chasanoff@gmail.com

Elnara Kazimova

Ganja, Azerbaijan

<https://orcid.org/0009-0006-8994-5351>

kazimovaelnara617@gmail.com

Naila Hasanova

Ganja, Azerbaijan

<https://orcid.org/0009-0003-1549-1690>

nailahesenova1974@gmail.com

Anar Guliyev

Ganja, Azerbaijan

<https://orcid.org/0009-0002-5549-4063>

anar_guliyev74@mail.ru

Artificial Intelligence and Machine Learning in Cybersecurity

Abstract

Artificial Intelligence (AI) and its subfield, Machine Learning (ML), are technologies that enable computers to think, learn, and make decisions like humans. The main algorithms of ML include supervised learning, unsupervised learning, and deep learning. Supervised learning works with labeled data to train models, while unsupervised learning identifies hidden patterns in unlabeled data. Deep learning, on the other hand, is based on multi-layered neural networks that mimic the human brain and solve complex problems using large datasets.

In the field of cybersecurity, AI and ML algorithms play critical roles in threat detection, identifying anomalies, and automating response systems. For example, deep learning algorithms can detect anomalies in network traffic and prevent DDoS attacks. Additionally, AI systems can reduce false positives, predict threats, and automate security measures.

However, the human factor remains irreplaceable in the application of AI in cybersecurity. Human expertise is crucial in areas such as strategic thinking, ethical decision-making, and creative problem-solving. In the future, collaboration between humans and AI in cybersecurity will increase, but challenges such as data quality, skill gaps, and ethical issues will also arise. Therefore, the future of cybersecurity will depend on the synergy between humans and AI.

This article examines the application, shortcomings, and advantages of AI and ML in various areas of cybersecurity, as well as the challenges ahead.

Keywords: information security, cybersecurity, cyberattacks, artificial intelligence, machine learning

Giriş

Süni intellekt - kompüterlərin insan kimi düşünmə, öyrənmə və qərar qəbul etmə bacarığını təmin edən texnologiyadır.

Maşın öyrənmə (Machine Learning) isə Sİ-nin bir alt sahəsi olub, verilənlərdən nümunələr çıxararaq modellərin təkmilləşdirilməsini təmin edir.

MÖ-nün Əsas Alqoritmləri

1. Nəzarətli Öyrənmə: Etiketlənmiş verilənlər üzərində öyrənmə.
2. Nəzarətsiz Öyrənmə: Verilənlərdə gizli strukturları aşkarlamaq.
3. Dərin Öyrənmə: Sinir şəbəkələrinə əsaslanan kompleks modellər.

Nəzarətli Öyrənmə Alqoritmı

Nəzarətli öyrənmə alqoritmı, verilənlərin giriş və çıxış nümunələri arasında əlaqə quraraq öyrənir. Bu alqoritm input (giriş) verilənləri ilə onların uyğun output (çıxış) dəyərləri üzərində təlim

keçirilir. Model öyrəndikdən sonra yeni input verilənləri üçün uyğun çıxışlar proqnozlaşdırıla bilər (Khakimov, 2023).

Tədqiqat

Əsas Mərhələlər:

Verilənlərin Toplanması: Keyfiyyətli və etiketlenmiş məlumatlar əldə edilir.

Modelin Təlimi: Alqoritm verilənləri analiz edərək nümunələri öyrənir.

Test və Təhlil: Model yeni verilənlər üzərində sınaq olaraq dəqiqlik qiymətləndirilir.

Populyar Nəzarətli Öyrənmə Alqoritmləri:

Qərar Ağları: Məlumatların iyerarxik struktura əsasən bölünməsi.

Logistik Regressiya: İki sinifli təsnifat problemləri üçün istifadə olunur.

SVM (Support Vector Machines): Məlumat nöqtələrini optimal hiper müstəviyə əsasən təsnif edir.

Misal olaraq, spam filtrləmə sistemləri nəzarətli öyrənmə metodundan istifadə edərək e-poçtların spam olub-olmadığını təsnif edir.

1. Nəzarətsiz Öyrənmə Alqritmi

Nəzarətsiz öyrənmə alqritmi, verilənlərdə gizli strukturların və nümunələrin müəyyənəşdirilməsi üçün istifadə olunur. Bu alqritmlər, giriş məlumatlarının etiketlenmədiyi hallarda tətbiq edilir və verilənlərdəki təbii qrupları və ya qruplaşmaları aşkarlamağa imkan verir (Namiot, Il'yushin, Chizhov, 2022).

Əsas Mərhələlər:

Verilənlərin Toplanması: Çoxsaylı xüsusiyyətlərdən ibarət məlumatlar toplanır.

Modelin Təlimi: Verilənlərdəki struktur və ya nümunələr analiz edilir.

Klasterləşdirmə və ya Azaldılma: Model məlumatları qruplaşdırır və ya ölçü azaldılmasını həyata keçirir.

Populyar Nəzarətsiz Öyrənmə Alqritmləri:

K-Ortalamalar Klasterləşməsi: Verilənləri müəyyən sayda qruplara ayırır.

Hierarxik Klasterləşmə: Verilənlərin iyerarxik struktura əsasən qruplaşdırılması.

Baş Komponent Analizi (PCA): Məlumatların ölçüsünü azaldaraq əsas komponentləri müəyyən edir.

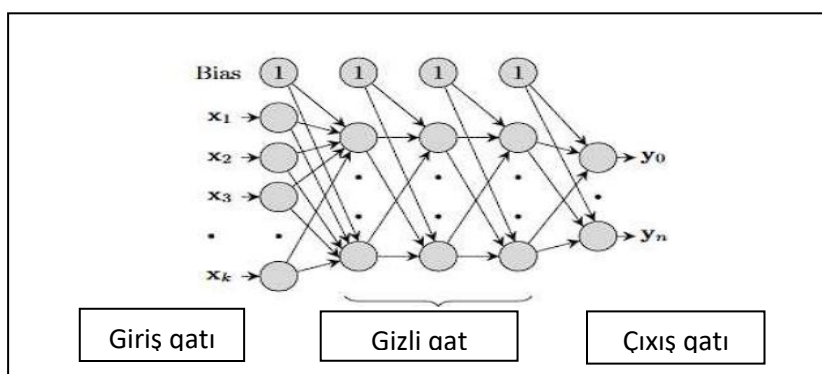
2. Dərin Öyrənmə Alqritmi (DL – Deep Learning)

Tətbiq Sahələri:

- Anomaliya Aşkarlanması: Şəbəkə trafikində normadan kənar nümunələrin müəyyən edilməsi.
- Müştəri Seqmentasiyası: İstehlakçıların davranışlarına görə qruplaşdırılması.
- Təhlükə Analizi: Kibertəhlükəsizlikdə naməlum zərərli proqramların aşkarlanması.

Misal üçün, şəbəkə monitorinqində nəzarətsiz öyrənmə anomal trafik nümunələrini aşkarlamaq üçün istifadə olunur. Bu, kibertəhlükəsizlik təhdidlərini vaxtında aşkarlamağa kömək edir.

Tədqiqatçılar DL-ni Şəkil 1-də göstərildiyi kimi ən azı üç gizli təbəqədən ibarət neyron şəbəkəsi kimi təyin edirlər (Roosmalen, Vranken, Eekelen, 2018). Tədqiqatçılar məlumatları dərin öyrənmə neyron şəbəkəsinin giriş qatında təmin edirlər. Daha sonra giriş təbəqəsi məlumatları mümkün çıxışları istehsal etmək üçün alqritmləri yerinə yetirən bir neçə gizli təbəqəyə göndərir.



Şəkil 1. Dərin öyrənmə neyron şəbəkəsi (Roosmalen, Vranken, Eekelen, 2018)

Dərin öyrənmə, insan beyninin fəaliyyətini təqlid edən çoxlaylı sinir şəbəkələrinə əsaslanır. Bu alqoritm böyük verilənlər dəstindən nümunələr öyrənərək mürəkkəb problemləri həll edə bilər.

Əsas Xüsusiyyətlər:

1. Çoxqatlı Sinir Şəbəkələri (Deep Neural Networks - DNNs): Dərin öyrənmə modelləri bir-birinə bağlı çoxsaylı gizli qatlara malikdir. Hər bir qat məlumatların fərqli xüsusiyyətlərini öyrənir.

2. Avtomatik Xüsusiyyət Çıxarışı: Ənənəvi alqoritmlərdən fərqli olaraq, dərin öyrənmə xüsusiyyətləri özü çıxarır.

3. Böyük Verilənlər: Dərin öyrənmə böyük miqdarda verilənlərdən istifadə edərək daha dəqiq nəticələr əldə edir.

Əsas Alqoritmlər

1. Konvolyusiya Sinir Şəbəkələri (Convolutional Neural Networks - CNNs): Şəkil emalı və obyekt tanıma üçün istifadə olunur.

2. Təkrar Sinir Şəbəkələri (Recurrent Neural Networks - RNNs): Vaxt ardıcılığına əsaslanan məlumatlar üçün, məsələn, mətn analizi və şəbəkə trafikinin monitorinqi.

3. Generativ Adversarial Şəbəkələr (GANs): Yeni verilənlər yaratmaq üçün istifadə olunur, məsələn, saxta məlumatların aşkarlanması.

Dərin Öyrənmənin Tətbiqi: Kibertəhlükəsizlikdə Anomaliya Aşkarlanması

Dərin öyrənmə alqoritmləri şəbəkə trafikində və ya istifadəçi davranışında anomaliyaları aşkar etmək üçün istifadə olunur (Grishin, Aver'yanova).

Misal: Dərin sinir şəbəkələri şəbəkə trafikinin təhlilində istifadə edilərək anomal nümunələri aşkar edə bilər. Bu metod xüsusilə DDoS hücumlarının vaxtında aşkarlanmasında effektivdir.

Kibertəhlükəsizlikdə Süni intellektin rolu

Süni intellekt artıq ənənəvi kibertəhlükəsizlik tədbirlərinin artırılmasında mühüm addımlar atıb. Təhlükənin aşkarlanmasından cavab avtomatlaşdırılmasına qədər, AI/ML alqoritmləri böyük həcmdə məlumatı insan imkanlarından çox-çox yüksək sürətlə təhlil edərək kibermüdafiə strategiyalarının səmərəliliyini və dəqiqliyini artırır. Bu alqoritmlər məlumatlarda nümunələri və anomaliyaları tanıya, proaktiv təhdidlərin aşkarlanmasına və cavab verməyə imkan verə bilər (Brown, 2022).

Kibertəhlükəsizlikdə AI-nin üstünlükləri

• **Təkmilləşdirilmiş təhlükənin aşkarlanması:** Süni intellekt sistemləri real vaxt rejimində anomaliyaları və potensial təhlükələri müəyyən etmək üçün böyük həcmdə məlumatları təhlil edə bilər. Ənənəvi üsullar çox vaxt müasir şəbəkələrin yaratdığı məlumatların böyük həcmi ilə mübarizə aparır. Süni intellekt alqoritmləri bu məlumatları daha sürətli emal edə bilər, məlum və naməlum təhlükələrin aşkarlanması sürətini yaxşılaşdırır.

• **Avtomatlaşdırılmış cavab:** Süni intellektin ən mühüm üstünlüklərindən biri aşkar edilmiş təhdidlərə cavabları avtomatlaşdırmaq qabiliyyətidir. Məsələn, potensial pozuntu müəyyən edildikdə, süni intellekt insan müdaxiləsi olmadan təsirlənmiş sistemləri təcrid edə və ya zərərli IP ünvanlarını bloklaya bilər. Bu sürətli cavab hücumçular üçün fürsət pəncərəsini minimuma endirir.

• **Azaldılmış yalançı pozitivlər:** Qabaqcıl AI alqoritmləri tarixi məlumat nümunələrindən öyrənməklə yanlış həyəcan siqnallarını minimuma endirmək üçün nəzərdə tutulmuşdur. Bu qabiliyyət kibertəhlükəsizlik mütəxəssisləri arasında xəbərdarlıq yorğunluğunu azaldır və onlara çoxsaylı yalan xəbərdarlıqları süzmək əvəzinə həqiqi təhlükələrə diqqət yetirməyə imkan verir.

• **Proqnoz analizi:** Süni intellekt istismardan əvvəl potensial zəiflikləri proqnozlaşdırmaq üçün tarixi məlumatlardan istifadə edə bilər. İstifadəçi davranışı və şəbəkə trafikindəki nümunələri təhlil edərək, AI sistemləri ransomware və ya fişinq kimi kiberhücumlarla əlaqəli riskləri proaktiv şəkildə azalda bilər (Smith, 2021).

Süni intellekt - şəbəkə trafikində, sistem loglarında və istifadəçi davranışında anomaliyaları aşkar etmək üçün müxtəlif maşın öyrənmə alqoritmlərindən istifadə edir. Bu alqoritmlər arasında **Qərar Ağları, Dəstək Vektor Maşınları (SVM) və Neuron Şəbəkələri** xüsusi yer tutur.

Anomaliya Aşkarlama

Anomaliya aşkarlama, normal davranışdan kənara çıxan nümunələri müəyyən etmək üçün istifadə olunur. Bu proses aşağıdakı düsturla modelləşdirilə bilər:

$$\text{Anomaliya Skoru} = \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2$$

Burada:

- x_i = Məlumat nöqtəsi
- μ = Məlumat dəstinin ortalaması
- n = Məlumat nöqtələrinin sayı

Yüksək anomaliya skoru, şəbəkədə potensial təhdidə işarə edir.

Təhlil Alqoritmlərinin Müqayisəsi

Aşağıdakı cədvəldə müxtəlif AI alqoritmlərinin anomaliya aşkarlamada effektivliyi müqayisə edilir:

Alqoritm	Dəqiqlik (%)	İşləmə Vaxtı (ms)
Qərar Ağacları	92	150
SVM	89	200
Neuron Şəbəkələri	95	300

Cədvəl 1: Müxtəlif AI alqoritmlərinin anomaliya aşkarlamada effektivliyi.

Süni İntellekt, xüsusilə **Təbii Dil Emalı (NLP)** alqoritmləri vasitəsilə, phishing e-poçtlarını və sosial mühəndislik hücumlarını aşkar edir. Bu alqoritmlər, mətnin dil quruluşunu, emosional tonunu və göndərənə reputasiyasını təhlil edir (Doe, 2020).

Phishing ehtimalı aşağıdakı loqistik reqressiya düsturu ilə hesablanır:

$$P(\text{Phishing}) = \frac{1}{1 + e^{-(w_1x_1 + w_2x_2 + \dots + w_nx_n)}}$$

Burada:

- w_i = Xüsusiyyət i -nin çəkisi
- x_i = Xüsusiyyət i -nin dəyəri

Bu təsirli imkanlara baxmayaraq, AI kibertəhlükəsizliyi tamamilə ələ keçirməyə hazır deyil. İnsan təcrübəsi bir neçə səbəbə görə əvəzolunmaz olaraq qalır (Goodfellow, Bengio, Courville, 2016):

- **Strateji düşüncə:** Süni intellekt nümunənin tanınması və sürətli təhlilində üstün olsa da, insan təhlükəsizliyi üzrə mütəxəssislərin gətirdiyi strateji düşüncə və kontekstual anlayışdan məhrumdur.

- **Etik mülahizələr:** Kibertəhlükəsizlik çox vaxt insan mühakiməsi və cavabdehlik tələb edən mürəkkəb etik qərarları əhatə edir.

- **Problemlərin həllində yaradıcılıq:** Kiber hücumçular daim yeni taktikalar inkişaf etdirirlər. İnsan yaradıcılığı yeni təhlükələri qabaqlamaq və onlara qarşı mübarizə aparmaq üçün vacibdir.
- **Nəticələrin təfsiri:** Süni intellekt anlayışlar yarada bilər, lakin insan ekspertləri bu tapıntıları təşkilatın təhlükəsizlik mövqeyinin daha geniş kontekstində şərh etməlidir.

Cədvəl 2. İnsan və Süni intellektin gələcək iş bölgəsi

İnsan komponenti	AI Komponenti
Strateji Planlaşdırma	Məlumat Təhlili
Etik mülahizələr	Pattern tanınması
Nəzarət və İdarəetmə	Avtomatlaşdırılmış Tapşırıqlar
Tənqidi Təfəkkür	Avtomatlaşdırma və Səmərəlilik
Problemlərin kompleks həlli	Tez cavab
Mədəni Şüur	Davranış Analitikası

Kibertəhlükəsizliyin gələcəyi insan intellektinin süni intellekt texnologiyaları ilə birlikdə işlədiyi əməkdaşlıq yanaşması əsasında olmalıdır. Cədvəl 1-dən də göründüyü kimi, kibertəhlükəsizlik mütəxəssisləri tənqidi düşünmə və problem həll etmə bacarıqlarını qoruyarkən süni intellekt alətlərindən səmərəli istifadə etmək üçün öz bacarıqlarını uyğunlaşdırmalıdır. Bu sinerji hər ikisinin güclü tərəflərini birləşdirir (Schneier, 2015).

Kibertəhlükəsizlikdə bu süni intellektlə genişləndirilmiş gələcəyi idarə edərkən bir neçə problem həll edilməlidir (Parisi, 2019):

- **Süni intellektlə işləyən hücumlar:** Müdafiəçilər süni intellektdən istifadə etdikləri kimi, hücumçular da bunu edə bilərlər. Bu silahlanma yarışını davamlı sayıqlıq və uyğunlaşma tələb edir.

- **Məlumat keyfiyyəti:** Süni intellekt sistemləri yalnız öyrədildikləri məlumatlar qədər yaxşıdır. Kibertəhlükəsizlikdə effektiv süni intellekt üçün yüksək keyfiyyətli, müxtəlif məlumat dəstlərinin təmin edilməsi çox vacibdir.

- **Bacarıq boşluğu:** Kibertəhlükəsizlik və süni intellekt arasında körpü yarada bilən, süni intellektlə idarə olunan təhlükəsizlik sistemlərini şərh edən və idarə edən mütəxəssislərə artan ehtiyac var.

- **Tənzimləyici uyğunluq:** Süni intellekt kibertəhlükəsizlikdə daha çox yayıldıqca, süni intellektdən istifadə və məlumatların qorunması ilə bağlı tənzimləyici mənzərədə naviqasiya getdikcə mürəkkəbləşir.

- **Texnologiyadan asılılıq:** Avtomatlaşdırılmış sistemlərə həddən artıq etibar etmək təşkilatların fundamental təhlükəsizlik təcrübələrini nəzərdən qaçırmalarına səbəb ola bilər.

Kibertəhlükəsizlikdə AI-nin gələcəyi çox ümidverici görünür. Bu sahənin bəzi əsas tendensiyaları və imkanları bunlardır (Chollet, 2018):

Avtomatlaşdırma və cavab sürəti:

Süni intellektə əsaslanan sistemlər insan müdaxiləsi olmadan təhdidləri dərhal aşkarlaya və onlara cavab verə biləcək, cavab müddətini əhəmiyyətli dərəcədə azaldacaq və uğurlu hücumlar ehtimalını azaldacaq.

Təhdid Kəşfiyyəti:

Süni intellekt böyük həcmdə məlumatları təhlil edə və ənənəvi üsullarla qaçırıla bilən mürəkkəb təhlükələri müəyyən edə biləcək. Maşın öyrənmə modelləri tarixi məlumatlara əsaslanaraq potensial hücumları proqnozlaşdırmağa biləcək.

Təkmilləşdirilmiş Aşkarlama Dəqiqliyi:

Müasir süni intellekt alqoritmləri artıq təhdidlərin aşkarlanmasında yüksək dəqiqlik nümayiş etdirir. Gələcəkdə bu alqoritmlər daha da dəqiqləşəcək və yanlış pozitivləri minimuma endirə biləcək.

Adaptiv sistemlərin təkamülü:

Süni intellektlə işləyən təhlükəsizlik sistemləri davamlı olaraq yeni təhdidləri öyrənəcək və onlara uyğunlaşacaq ki, bu da onlara təcavüzkarlardan bir addım öndə olmağa imkan verəcək.

İnsan-AI əməkdaşlığı:

Süni intellekt kibertəhlükəsizliyin ayrılmaz hissəsinə çevrilsə də, insan amili vacib olaraq qalacaq. Kibertəhlükəsizlik mütəxəssisləri strateji düşüncəni və yaradıcılığı qoruyaraq öz işlərini təkmilləşdirmək üçün süni intellektdən bir vasitə kimi istifadə edəcəklər.

Etika və Məxfilik:

Kibertəhlükəsizlikdə süni intellektin gələcəyinin vacib aspektlərindən biri süni intellekt texnologiyalarından etik və məxfi istifadənin təmin edilməsi olacaq. Standartların və qaydaların hazırlanması və tətbiqi əsas olacaqdır (MITRE Corporation, 2021).

AI hücumçularına qarşı mübarizə:

Təhlükəsizlik texnologiyalarının inkişafına cavab olaraq, təcavüzkarlar da hücumları həyata keçirmək üçün süni intellektdən istifadə etməyə başlayacaqlar. Belə AI hücumçularına qarşı mübarizə kibertəhlükəsizlik mütəxəssisləri üçün yeni problem olacaq.

Beləliklə, kibertəhlükəsizlikdə AI-nin gələcəyi əhəmiyyətli irəliləyişlər və təkmilləşdirmələr vəd edir, eyni zamanda yeni problemlərin və problemlərin həllini tələb edəcəkdir. İnsanların və süni intellektin səylərini birləşdirmək kibertəhlükələrə qarşı uğurlu müdafiənin açarı olacaq (European Union Agency for Cybersecurity (ENISA), 2022).

Nəticə

Kibertəhlükəsizlik sahəsində süni intellekt və maşın öyrənmə texnologiyaları mühüm inqilabi dəyişikliklər edərək, təhdidlərin daha effektiv aşkarlanması və qarşısının alınması imkanı yaradır. Bu texnologiyalar, müasir kibertəhlükəsizlik təhdidlərinə qarşı daha effektiv müdafiə tədbirlərinin inkişaf etdirilməsinə kömək edir. Lakin, süni intellekt texnologiyalarının tam potensialını əldə etmək üçün insan intellekti ilə əməkdaşlıq və bəzi problemlərin həll edilməsi vacibdir. Süni intellektin kibertəhlükəsizlikdəki rolu və potensialı, bu sahənin gələcəyinin parlaq olacağını göstərir. Sİ-in kvant hesablama, blokçeyn və 5G texnologiyaları ilə integrasiyası, kibertəhlükəsizliyin daha da inkişaf etdirilməsi üçün geniş imkanlar yaradır. Bundan əlavə, İzah Edilə Bilən AI (XAI) inkişafı, sistemlərin şəffaflığını artıracaq.

Ədəbiyyat

1. Brown, K. (2022). *Artificial Intelligence in Cybersecurity*. Springer.
2. Chollet, F. (2018). *Deep Learning with Python*. Manning Publications.
3. Doe, A. (2020). *Advanced Cybersecurity Techniques*. Oxford University Press.
4. European Union Agency for Cybersecurity (ENISA). (2022). *Artificial Intelligence in Cybersecurity*. <https://www.enisa.europa.eu>
5. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
6. Grishin, I. O., Aver'yanova, A. N. *Primeneniye iskusstvennogo intellekta v kiberbezopasnosti. Zhurnal «Nauchnyy lider»* vypusk # 49 (199).
7. Khakimov, A. A. (2023, Noyabr). *Rol' iskusstvennogo intellekta v kiberbezopasnosti. Nauchnaya stat'ya, opublikovannaya v zhurnale Universum: tekhnicheskoye nauki*. №11 (116).
8. MITRE Corporation. (2021). *Adversarial Threat Landscape for AI Systems*. <https://www.mitre.org> saytı.
9. Namiot, D. Ye., Il'yushin, Ye. A., Chizhov, I. V. (2022). *Iskusstvennyy intellekt i kiberbezopasnost'. International Journal of Open Information Technologies* ISSN: 2307-8162 vol. 10, № 9.
10. Smith, J. (2021). *Machine Learning for Cyber Threat Detection*. Wiley.

11. Schneier, B. (2015). *Data and Goliath: The Hidden Battles to Collect Your Data and Control Your World*. W.W. Norton & Company.
12. Parisi, A. (2019). *Hands-On Artificial Intelligence for Cybersecurity: Implement smart AI systems for preventing cyber-attacks and detecting threats and network anomalies*. Packt Publishing.

Daxil oldu: 14.11.2024

Qəbul edildi: 13.02.2025