

UOT 004.738; 004.82; 004.93

**KİBERTƏHLÜKƏSİZLİYİN SÜNİ İNTELLEKT ALQORİTMLƏRİ VƏ BİG DATA
ANALİTİKASI İLƏ GÜCLƏNDİRİLMƏSİ: ÇƏTİNLİKLƏR VƏ HƏLLƏR****Elvin Amil oğlu Muradzadə**elvin.muradzada@mdu.edu.az

Mingəçevir Dövlət Universiteti

<https://doi.org/10.30546/SJSD.2025.5.1.049>

Xülasə: Rəqəmsal transformasiya sürətləndikcə, kibertəhlükəsizlik mürəkkəb təhdidlərlə üzləşir. Bu tədqiqat süni intellekt alqoritmləri və böyük verilənlər analitikasının kibertəhlükəsizliyi gücləndirmək üçün integrasiyasını və şəffaflıq problemlərinə fokuslanaraq araşdırılır. Təsviri-keyfiyyət metodologiyası ilə ədəbiyyat, nümunələr və ikincil verilənlər təhlil edilərək effektivlik qiymətləndirilib. Ümumilikdə dərin öyrənmə 93 % dəqiqliklə anomaliya aşkarlayır, bu öz növbəsində qayda əsaslı metodlardan 18 % üstündür. Nümunə olaraq Spark 500 GB veriləni 35 saniyədə emal edir (Hadoop: 60 saniyə). Dövri yeniləmələr dəqiqliyi 15 % artırır. Süni intellekt şəffaflığı və integrasiya çətinliklərini real vaxtda aşkarlamayı çətinləşdirir. Bu tədqiqat mövcud sahəyə sistem dayanıqlığını artırmaq üçün strategiyalar təklif etməklə töhfə verir; bu strategiyalara şəffaflıq üçün izah edilə bilən süni intellektin qəbulu və qabaqcıl verilənlər integrasiya texnikaları daxildir. Gələcək tədqiqatlar sektorlara xas təbiiqlərə və cavab mexanizmlərinə yönəlməlidir.

Açar sözlər: kibertəhlükəsizliyin optimallaşdırılması, süni intellekt (Sİ), BİG Data təhlili, təhdidlərin aşkarlanması

Giriş

Rəqəmsallaşma müxtəlif sektorlarda artdıqca, kibertəhlükəsizlik əsas prioritetə çevrilmişdir. İnformasiya texnologiyalarındakı sürətli inkişaf insan həyatının bir çox sahəsində əhəmiyyətli irəliləyişlərə səbəb olsa da, eyni zamanda daha mürəkkəb və müxtəlif kiber hücum risklərini də artırmışdır [1]. Kiber hücumlar artıq kiçik miqyaslı cinayətkar qruplarla məhdudlaşmır, ölkələrin kritik infrastrukturunu – səhiyyə, energetika, maliyyə və dövlət sektorlarını – poza biləcək təhdidlərə çevrilmişdir. Cybersecurity Ventures-ın hesabatına görə, 2025-ci ilin sonlarına qədər kiber cinayətlərin global xərclərinin 10,5 trilyon dollara çatacağı proqnozlaşdırılır ki, bu da 2015-ci ildəki 3 trilyon dollardan əhəmiyyətli artımı göstərir və kibertəhdidlərin iqtisadi təsirləri ilə yanaşı, global təhlükəsizlik və sabitlik üçün nəticələrini vurğulayır [2].

Kibertəhdidlərin maddi təsirinin diqqətəlayiq nümunələrindən biri 2021-ci ildə ABŞ-da səhiyyə sektoruna qarşı həyata keçirilən fidyə proqramı hücumudur. Bu hücum əhəmiyyətli maliyyə itkilərinə səbəb olmuş və kritik tibbi məlumatlara girişi kəsərək xəstə təhlükəsizliyini təhdid etmişdir. Fidyə proqramı səhiyyə şəbəkələrini yoluxduraraq, xəstə məlumatlarına əlçatanlığı kilidləmiş və girişin bərpası üçün ödəniş tələb etmişdir [3]. Bu hadisə mövcud kibertəhlükəsizlik strategiyalarının məhdudiyyətlərini vurğulayır ki, bu strategiyalar yeni yaranan təhdidlərə qarşı həssas infrastruktur qorumaqda getdikcə qeyri-kafi qalır. Süni intellekt (Sİ) və böyük verilənlər analitikası daha adaptiv və cavabdeh imkanlar təklif edərək kibertəhlükəsizlik dayanıqlığını gücləndirə bilər. Süni intellekt alqoritmləri təhdidlərin avtomatlaşdırılmış aşkarlanmasını və cavabını təmin edərək təşkilatlara hücumları real vaxtda aradan qaldırmağa və təhdid identifikasiyasının dəqiqliyini artırmağa imkan verir [4]. Məsələn, maşın öyrənmə alqoritmləri böyük verilənlər dəstlərini təhlil edərək potensial kibertəhdidləri işarə edən nümunələri və ya anomaliyaları aşkarlaya bilər. Süni intellekt şəbəkədə qeyri-adi fəaliyyət və ya şübhəli cəhdlər kimi anormallıqları daha sürətlə müəyyən edir, təşkilatlara bu təhdidlər böyük hücumlara çevrilməzdən əvvəl reaksiya verməyə imkan yaradır.

Böyük verilənlər analitikası şəbəkə fəaliyyət qeydləri, istifadəçi məlumatları və IoT cihazlarından alınan verilənlər kimi çoxsaylı mənbələrdən böyük həcmli verilənləri idarə etmək və təhlil etmək qabiliyyətini təklif edir. Bu verilənlər potensial təhdidləri göstərən nümunələri müəyyən

etməyə və keçmiş verilənlərə əsaslanaraq mümkün hücumları proqnozlaşdırmağa kömək edə bilər. Böyük verilənlər təhlili həmçinin proqnoz dəqiqliyini artırmağa və təhdidlərin identifikasiyasını sürətləndirməyə töhfə verərək daha sürətli cavab müddətini təmin edir. Süni intellekt və böyük verilənlərin integrasiyası ilə kibertəhlükəsizlik sistemləri yalnız davam edən təhdidləri aşkar etmir, həm də hücumları baş verməzdən əvvəl izləyir və qarşısını alır. Lakin süni intellekt və böyük verilənlər kibertəhlükəsizliyi gücləndirmək üçün böyük potensiala malik olsa da, onların tətbiqi ciddi çətinliklərlə üzləşir. Əsas çətinliklərdən biri, xüsusilə dərin öyrənmə əsaslı süni intellekt alqoritmlərinin şəffaflığıdır. Dərin öyrənmə texnikalarına əsaslanan bir çox süni intellekt alqoritmləri qərar qəbul etmə proseslərinin mürəkkəbliyi səbəbindən "qara qutu" adlandırılır və bu, insanların anlaması və şərh etməsi üçün çətinlik yaradır. Bu, qərarların hesabatlı və şəffaf olmasının zəruri olduğu kibertəhlükəsizlik kontekstində, xüsusilə maliyyə və səhiyyə kimi həssas sektorlarda problem yaradır.

Ədəbiyyat icmalı

Süni intellekt (SI) alqoritmləri problemləri həll etmək və ya tapşırıqları insan müdaxiləsi olmadan avtomatik yerinə yetirmək üçün nəzərdə tutulmuş addımlar və ya təlimatlar silsiləsidir. Bu alqoritmlər insan qabiliyyətlərini təqlid edir və verilənləri giriş kimi istifadə edərək qərar qəbulunu təmin edir və istənilən nəticələri əldə edir. Ənənəvi proqramlaşdırma metodlarından fərqli olaraq, əvvəlcədən müəyyən edilmiş qaydalara əsaslanmayan süni intellekt alqoritmləri, işlətdikləri verilənlərdən öyrənərək öz performanslarını təkmilləşdirə bilər. Bu texnologiya səs tanıma və verilənlər təhlili kimi müasir yeniliklərin əsasını təşkil edir. Süni intellektin əsas sahələrindən biri olan maşın öyrənməsində alqoritmlər proqnozlar və ya qərarlar qəbul edə bilən modelləri öyrətmək üçün istifadə olunur; nəticələr giriş verilənlərindən əldə edilir. Məsələn, nəzarətli öyrənmə alqoritmləri etikətlənmiş verilənlər dəstlərindən istifadə edərək gələcək nəticələri proqnozlaşdırır, nəzarətsiz öyrənmə isə etikətlənməmiş verilənlərdə nümunələri aşkar edir. Bu alqoritmlər məhsul tövsiyələri, saxtakarlıq aşkarlanması və görüntü klasterləşdirilməsi kimi müxtəlif tətbiqlərdə mühüm rol oynayır [2].

Süni intellekt (SI) alqoritmlərinin geniş istifadə olunan nümunələrindən biri süni neyron şəbəkələridir. Bu alqoritmlər insan beyinin funksiyalarını təqlid edərək, bir-biri ilə əlaqəli emal vahidlərindən istifadə edir və verilənləri təhlil edib işləyir. Simulyasiya edilmiş neyron şəbəkələri (SNN) üz tanıma, təbii dil emalı və səhiyyə diaqnostikası kimi sahələrdə tətbiq olunur. Daha mürəkkəb versiyalar, məsələn, dərin öyrənmə, böyük həcmli və mürəkkəb strukturlu verilənlərin emalına imkan verir ki, bu da onları görüntü və video təhlili üçün xüsusilə faydalı edir. Süni intellekt alqoritmləri həmçinin avtonom qərar qəbulunda mühüm rol oynayır. Məsələn, avtonom nəqliyyat vasitələri sensor verilənlərini təhlil edərək təhlükəsiz naviqasiya, maneə aşkarlanması və sürət nəzarəti ilə bağlı anı qərarlar qəbul etmək üçün süni intellekt alqoritmlərindən istifadə edir. Bu alqoritmlər təhlükəsizliyi təmin etmək üçün sürətli və dəqiq işləməlidir ki, bu da süni intellekt tətbiqlərində sürət və dəqiqliyin əhəmiyyətini vurğulayır.

Üstünlüklərinə baxmayaraq, süni intellekt alqoritmlərinin tətbiqi şəffaflıq və etika sahəsində çətinliklər yaradır. Xüsusilə dərin öyrənmə modellərindən istifadə edən bəzi süni intellekt alqoritmləri qərar qəbul proseslərinin mürəkkəbliyi səbəbindən tez-tez "qara qutu" adlandırılır və bu prosesləri ətraflı izah etmək çətindir. Bu, xüsusilə səhiyyə, hüquq və ya maliyyə sahələrində tətbiqlərdə qərarların necə qəbul edildiyi ilə bağlı narahatlıqlar doğurur. Nəticədə şəffaf, şərh edilə bilən və etik standartlara uyğun süni intellekt alqoritmlərinin yaradılması gələcək üçün əhəmiyyətli bir çağırış olaraq qalır [3].

Böyük verilənlər analitikası böyük həcmli verilənlərin toplanması, emalı və təhlilini əhatə edən çoxşaxəli prosesdir ki, bu da mənalı nəticələr əldə etməyə xidmət edir. Verilənlərin böyük həcmi və müxtəlifliyi səbəbindən bu təhlil ənənəvi metodların idarə edə bilmədiyi verilənləri emal etmək üçün nəzərdə tutulmuş qabaqcıl texnologiyalardan istifadə edir. Təhlil olunan verilənlər sosial media, biznes əməliyyatları, sensorlar və əşyaların interneti (IoT) cihazları kimi müxtəlif mənbələrdən əldə edilə bilər. Buna görə də, böyük verilənlər analitikası təşkilatlar üçün verilənlərdəki nümunələri, tendensiyaları və gizli əlaqələri aşkar etməkdə mühüm rol oynayır. Böyük

verilənlər analitikasının əsas yanaşmalarından biri maşın öyrənməsidir ki, bu, kompüter sistemlərinə açıq proqramlaşdırma olmadan verilənlərdən öyrənməyə imkan verir. Maşın öyrənmə alqoritmləri verilənlərdəki nümunələri tanıyır və keçmiş məlumatlara əsaslanaraq proqnozlar verir. Bundan əlavə, statistik texnikalar və süni intellekt daha dərin nəticələr əldə etmək üçün istifadə olunur, bu da biznes və tədqiqat kontekstlərində qərar qəbulunu dəstəkləyən daha dəqiq və aktual analitik nəticələr təmin edir. Verilənlərin emal sürəti böyük verilənlər analitikasının mühüm elementinə çevrilmişdir. Real vaxt rejimində təhlil sayəsində verilənlər toplandıqdan dərhal sonra emal edilə və qiymətləndirilə bilər, bu da vaxtında qərar qəbulunu dəstəkləyən anı nəticələr verir. Bu, xüsusilə maliyyə, səhiyyə və təhlükəsizlik kimi sənayelərdə vacibdir, çünki burada qərarların dəqiqliyi və sürəti nəticələri müəyyən etməkdə kritik rol oynayır. Məsələn, kibertəhlükəsizlikdə real vaxt analizi təhdidləri tez aşkar edərək, daha böyük zərər baş verməzdən əvvəl qabaqlayıcı tədbirlər görməyə imkan.

Kibertəhlükəsizlik sistemlərinin optimizasiyası

Kibertəhlükəsizlik sistemlərinin gücləndirilməsi, getdikcə mürəkkəbləşən rəqəmsal infrastruktur müxtəlif təhdidlərdən qorumaq üçün mühüm addımdır. Texnologiya inkişaf etdikcə, hakerlik, zərərli proqramlar və məlumat oğurluğu kimi kiber hücumların tezliyi əhəmiyyətli dərəcədə artmışdır. Buna qarşı kibertəhlükəsizliyin optimizasiyası təşkilatların müdafiəsini kibertəhdidlərin müəyyənləşdirilməsi, zərərsizləşdirilməsi və aradan qaldırılmasında möhkəmləndirməyi hədəfləyir. Bu proses təkcə texnoloji yeniləmələri deyil, həm də təhlükəsizlik siyasətlərinin gücləndirilməsini, istifadəçilərin təlimatlandırılmasını və daha sərt protokolların tətbiqini əhatə edir.

Təhlükəsizliyin optimizasiyasında geniş istifadə olunan yanaşmalardan biri daha qabaqcıl təhdid aşkarlama texnologiyalarının qəbuludur. Süni intellekt və maşın öyrənməsi kimi texnologiyalar sistemlərə şəbəkələrdə potensial kiber hücumları işarə edə biləcək qeyri-adi nümunələri müəyyən etmək imkanı verir. Real vaxt rejimində təhlil sayəsində bu sistemlər hücumlar sistemə təsir etməzdən əvvəl daha tez xəbərdarlıq edə bilər. Bundan əlavə, güclü şifrələmə həssas verilənləri icazəsiz girişdən qoruyur. Texnologiyadan kənarda, kibertəhlükəsizliyin optimizasiyası təşkilat daxili siyasətlərin möhkəmləndirilməsini də tələb edir. Hər bir təşkilat bütün işçilərin təhlükəsizliyin saxlanması əhəmiyyətini anlamasını təmin etmək üçün aydın və strukturlaşdırılmış təhlükəsizlik siyasətləri hazırlamalıdır. Kiber təhlükəsizlik riskləri və qabaqlayıcı tədbirlər, məsələn, fişinq e-poçtlarını tanıma və parolları təhlükəsiz saxlama üzrə müntəzəm təlimlər bu siyasətlərin vacib hissəsidir. Siyasətlərin ardıcıl tətbiqi ilə istifadəçi diqqətsizliyi və ya məlumatlılıq çatışmazlığı nəticəsində yaranan hücum riskləri azaldıla bilər. Çoxqatlı təhlükəsizlik kibertəhlükəsizlik optimizasiyasının digər əsas elementidir. Bu strategiya sistem daxilində bir neçə qoruma qatını təmin edir; təhlükəsizlik divarları, antivirus proqramları və müdaxilə aşkarlama sistemləri kimi texnologiyalar kiber hücumların qarşısını almaq üçün birgə fəaliyyət göstərir [4].

Süni intellekt alqoritmlərinin integrasiyasının kibertəhlükəsizliyin gücləndirilməsinə təsiri

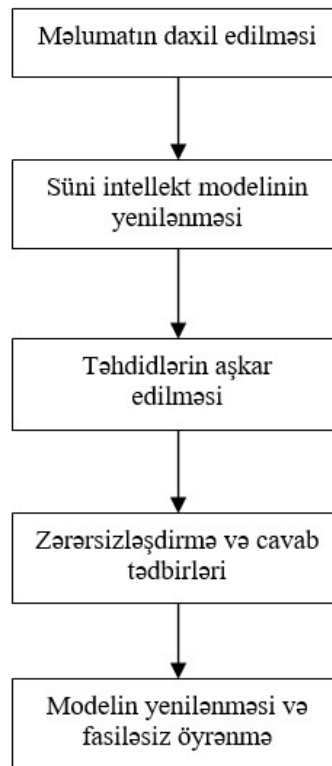
Süni intellekt alqoritmlərinin böyük verilənlər analitikası ilə integrasiyası kibertəhlükəsizlik sistemlərinin fəaliyyətini əhəmiyyətli dərəcədə təkmilləşdirmişdir. Bu sinerjinin əsas nəticələrindən biri təhdidlərin aşkarlanma qabiliyyətinin artmasıdır. Maşın öyrənmə alqoritmlərindən istifadə edərək, sistemlər böyük həcmli verilənləri real vaxt rejimində emal edərək hücumu işarə edə biləcək nümunələri və anomaliaları müəyyən edə bilər. Bu qabiliyyət kibertəhdidlərə daha sürətli cavab verməyə imkan yaradır və beləliklə, təhlükəsizlik pozuntularından yaranan potensial itkiləri azaldır. Bundan əlavə, bu integrasiya kibertəhlükəsizlik idarəetməsində avtomatlaşdırmanı asanlaşdırır. Süni intellekt alqoritmləri verilənləri avtomatik izləyir və təhlil edir, adi tapşırıqlar üçün insan müdaxiləsinə ehtiyacı azaldır; bu proseslər çox vaxt tələb edir və səhvlərə meyllidir. Avtomatlaşdırma ilə təhlükəsizlik qrupları strateji və dərin təhlilə daha çox diqqət yetirə bilər, alqoritmlər isə ilkin təhdid aşkarlamasını və cavabını daha yüksək səmərəliliklə idarə edir. Bu integrasiyanın digər üstünlüyü hücumları baş verməzdən əvvəl proqnozlaşdırmaq və zərərsizləşdirmək qabiliyyətidir. Böyük verilənlər analitikası vasitəsilə kibertəhlükəsizlik sistemləri şübhəli nümunələri və davranışları müəyyən edərək potensial təhdidlərə qarşı lazım olan məlumatları təmin edir. Məsələn,

sistemlər keçmiş hücumların verilənlərini təhlil edərək zəiflikləri müəyyən edən və hücum baş verməzdən əvvəl müvafiq qoruyucu tədbirlər tətbiq edən proqnoz modelləri hazırlaya bilər. Lakin süni intellekt alqoritmlərinin kibertəhlükəsizlikdə tətbiqi öz çətinliklərini də yaradır. Süni intellekt təhdidlərin aşkarlanması və qarşısının alınmasında kömək edə bilsə də, hücumçular bu texnologiyadan daha mürəkkəb hücum metodları hazırlamaq üçün sui-istifadə edə bilərlər. Məsələn, hücumçular süni intellekt alqoritmlərindən istifadə edərək hücumları avtomatlaşdırma və sistem zəifliklərindən istifadə edə bilərlər. Buna görə də, təşkilatların təhlükəsizlik strategiyalarını davamlı olaraq yeniləməsi və təkmilləşdirməsi effektivlik və müasirlik baxımından vacibdir [5].

Kibertəhlükəsizlik optimizasiyası üçün sistemlərin inteqrasiyasındakı çətinliklər

Süni intellekt alqoritmlərinin böyük verilənlər analitikası ilə inteqrasiyası kibertəhlükəsizlik sistemlərinin gücləndirilməsində bir neçə əsas çətinlik yaradır. Birincisi, kibertəhlükəsizlikdə istifadə olunan verilənlər yüksək dərəcədə mürəkkəbdir. Fəaliyyət qeydləri, şəbəkə trafiki və istifadəçi qarşılıqlı əlaqələri kimi müxtəlif mənbələrdən toplanan məlumatlar böyük həcmə və əhəmiyyətli dəyişkənliyə malikdir. Bu verilənləri real vaxt rejimində təhdid aşkarlanması üçün effektiv təhlil etmək və emal etmək üçün həm sürətli, həm də müxtəlif verilən tiplərini idarə edə bilən alqoritmlər tələb olunur (Admass və s., 2024).

İkinci çətinlik süni intellekt modellərinin yüksək keyfiyyətli və reprezentativ verilənlər dəstləri üzərində təlimi ilə bağlıdır. Çox vaxt təlim üçün istifadə olunan verilənlər real kibertəhlükəsizlik risklərini dəqiq əks etdirmir. Kibertəhdidlər daim dəyişir, bu da köhnəlmiş və ya qeyri-kafi verilənlərlə hazırlanmış modellərin yeni hücum tiplərini effektiv şəkildə müəyyən edə bilməməsinə səbəb olur. Buna görə də, süni intellekt sistemlərinin təhdidləri effektiv müəyyən edib zərərsizləşdirməsi üçün müvafiq verilənlər dəstlərinin toplanması və saxlanması, düzgün təlim və optimizasiya vacibdir. Süni intellekt və böyük verilənlər inteqrasiyasında digər çətinlik şəffaflıqla bağlıdır.



Şək 1. Süni intellekt sistemlərində dövri məlumat yeniləmələri ilə təhdid aşkaretmə axını

Xüsusilə dərin öyrənmə əsaslı süni intellekt alqoritmləri qərar qəbul proseslərinin mürəkkəbliyi səbəbindən tez-tez “qara qutu” adlandırılır və bunların şərh edilməsi çətinidir. Kibertəhlükəsizlik kontekstində qərarların əsaslandırılması zəruri olduğundan, bu şəffaflıq çatışmazlığı texnologiyanın geniş qəbuluna mane ola bilər. Təşkilatlar istifadəçilər və maraqlı tərəflər arasında etimadı artırmaq üçün süni intellekt modellərinin qərarlarının əsaslarını izah edə biləcək üsullar hazırlamalıdır. Verilənlərin məxfiliyi və qanunauyğunluq tələbləri ilə bağlı məsələlər əlavə çətinliklər yaradır. Böyük verilənlərin kibertəhlükəsizlikdə istifadəsi artdıqca, təşkilatlar həssas məlumatların idarə edilməsi və qorunmasında diqqətli olmalıdır. Süni intellekt inteqrasiyası verilənlər pozuntusu riskini artırır, bu da şirkətin nüfuzuna zərər vura və hüquqi cəzalara səbəb ola bilər. Buna görə bütün verilənlərin toplanması və təhlili fəaliyyətlərinin GDPR kimi müvafiq məxfilik qanunlarına uyğunluğunu təmin etmək vacibdir.

Nəhayət, başqa bir mühüm çətinlik fənlərarası əməkdaşlıq ehtiyacıdır. Süni intellekt və verilənlər inteqrasiyası vasitəsilə kibertəhlükəsizlik sistemlərinin optimizasiyası informasiya texnologiyaları, kibertəhlükəsizlik və verilənlər analitikası üzrə mütəxəssislərin əməkdaşlığını tələb edir. Bu müxtəlif qruplar çətinlikləri dərinlən anlamalı və texnologiyanın onları necə həll edə biləcəyini bilməlidir. Effektiv əməkdaşlıq sayəsində təşkilatlar yeni yaranan kibertəhdidlərə qarşı hazırlığını gücləndirərək təhlükəsizlik strategiyalarının aktuallığını və effektivliyini qoruya bilər.

Şəx. 1 müntəzəm məlumat yeniləmələrindən istifadə etməklə, təhdid aşkar etmə axını nümayiş etdirir və süni intellekt modellərinin dövrü yenilənməsinin yeni yaranan təhdidlərin qarşısının alınmasına necə kömək etdiyini təsvir edir.

Bu tədqiqatın nəticələri göstərir ki, süni intellektin və geniş miqyaslı məlumat analizinin kibertəhlükəsizlik sahəsində əhəmiyyətli potensiala malik olmasına baxmayaraq, bəzi sahələr, xüsusilə, süni intellektin izah edilə bilməsi və məlumatların inteqrasiyası, əlavə inkişaf tələb edir. Həmçinin, əldə olunan nəticələr sübut edir ki, kibertəhlükəsizlik sistemlərinin yeni yaranan təhdidlərə qarşı effektivliyini və aktuallığını təmin etmək üçün təlim məlumatlarının mütləq yenilənməsi vacibdir.

Tədqiqat metodu

Bu tədqiqat kibertəhlükəsizlik sistemlərinin möhkəmliyini artırmaqda süni intellekt və verilənlər analitikasının potensial istifadəsini araşdırmaq və müəyyənləşdirmək üçün təsvir-keyfiyyət metodundan istifadə edir. Bu yanaşma kibertəhdidlərin müəyyənləşdirilməsi və idarə edilməsində tətbiq olunan metodların dərinlən anlaşılmasını təmin etmək, həmçinin süni intellekt və böyük verilənlər texnologiyalarının real dünya mühitlərində tətbiqi ilə üzləşən çətinlikləri öyrənmək üçün seçilmişdir. Tədqiqatın təhlili üçün nümunə araşdırması dizaynı qəbul edilmişdir və müxtəlif ilkin və ikincil mənbələrdən, məsələn, kibertəhlükəsizlik insident hesabatları, ədəbiyyat icmalları və müvafiq əvvəlki tədqiqatlardan alınan verilənlərdən istifadə olunmuşdur. Verilənlərin toplama prosesi ədəbiyyat icmalı və ikincil mənbələr vasitəsilə həyata keçirilmişdir.

Ədəbiyyat icmalı kibertəhlükəsizlik, süni intellekt və böyük miqyaslı verilənlər analitikası sahələrində akademik nəşrləri, elmi jurnalları və tədqiqat hesabatlarını əhatə edərək kibertəhdid aşkarlamasında istifadə olunan süni intellekt və böyük verilənlər analitikası texnikalarını müəyyənləşdirmişdir. Bundan əlavə, ikincil verilənlər kibertəhlükəsizlik şirkətlərinin illik hesabatlarından, qlobal statistik məlumatlardan və müxtəlif sektorlarda kiber hücumların nümunə araşdırmalarından toplanmışdır. Bu verilənlər hücum nümunələrini və texnologiyaların tətbiqi çətinliklərini anlamaq üçün istifadə edilmişdir.

Verilənlərin təhlili tematik analiz yanaşması ilə aparılmış, nəticələrdən ortaya çıxan nümunələr və əsas mövzular müəyyənləşdirilmişdir. Proses verilənlərin üç əsas kateqoriyaya təşkil edilməsi və kodlaşdırılması ilə başlamışdır:

- 1) süni intellekt alqoritm növləri;
- 2) böyük verilənlər analitikası metodları;
- 3) tətbiq ilə bağlı çətinliklər və imkanlar.

Birinci kateqoriyada dərin öyrənmə və SNN kimi süni intellekt alqoritmləri ilə bağlı verilənlər anomaliyaların aşkarlanma dəqiqliyi və tətbiq oluna bilməsi baxımından təhlil edilmişdir.

İkinci kateqoriya Hadoop və Spark kimi böyük verilənlər platformalarına fokuslanaraq, onların sürətini və böyük miqyaslı kibertəhlükəsizlik verilənlərinin idarə edilməsində integrasiya çevikliyini qiymətləndirmişdir.

Üçüncü kateqoriya süni intellekt şəffaflığı və verilənlərin integrasiyası problemləri daxil olmaqla çətinlikləri, optimizasiya imkanları ilə birlikdə müəyyənləşdirmişdir [6]. Nəticələr tədqiqat məqsədləri kontekstində şərh olunmuş, xüsusilə süni intellekt və böyük verilənlərin kibertəhlükəsizlik dayanıqlığını artırmaq üçün necə optimallaşdırıla biləcəyi vurğulanmışdır. Bundan əlavə, nəticələr əvvəlki tədqiqatlarla çarpaz doğrulama aparılaraq boşluqlar və əlavə araşdırma tələb edən sahələr, məsələn, izah edilə bilən süni intellekt və adaptiv verilənlər integrasiya sistemlərinə ehtiyac vurğulanmışdır. Şəkil 1 bu tədqiqatın çərçivəsini göstərir.

Nəticə

Bu tədqiqat süni intellekt (Sİ) və böyük miqyaslı verilənlər təhlilinin kibertəhlükəsizlik sistemlərini gücləndirmək, xüsusilə kibertəhdidlərin aşkarlanması və qarşısının alınması baxımından əhəmiyyətli potensial təklif etdiyi qənaətinə gəlir. SNN-lər və dərin öyrənmə kimi texnikalar kiber hücumlarda mürəkkəb anomaliyaların və nümunələrin aşkarlanmasında yüksək dəqiqlik göstərmişdir. Bundan əlavə, böyük miqyaslı verilənlər təhlili geniş verilənlər dəstlərinin real vaxt rejimində emalına imkan verərək, təhdid aşkarlama sürətini artırmışdır. Lakin tədqiqat süni intellektin şəffaflığı və böyük miqyaslı verilənlər təhlilində verilənlərin integrasiya çevikliyi kimi həll edilməli olan bir neçə çətinliyi də müəyyənləşdirir. Süni intellekt alqoritmlərinin şəffaflıq məhdudiyyətləri yüksək hesabatlılıq tələb edən tətbiqlərdə maneə yarada bilər, yüksək verilən mürəkkəbliyi isə səmərəli təhlil üçün təkmilləşdirilmiş integrasiya sistemlərini tələb edir. Bu nəticələrə əsasən, süni intellekt və böyük miqyaslı verilənlər təhlilinin kibertəhlükəsizlikdə effektivliyini artırmaq üçün bir neçə tövsiyə təklif olunur.

Birincisi, alqoritmik qərar qəbul proseslərinin istifadəçilər tərəfindən anlaşılmasını yaxşılaşdırmaq üçün, xüsusilə yüksək şəffaflıq tələb edən kontekstlərdə, izah edilə bilən süni intellekt kimi daha şərh edilə bilən süni intellekt texnologiyalarının inkişafı vacibdir.

İkincisi, təşkilatlara müxtəlif mənbələrdən gələn müxtəlif verilən formatlarını idarə edə bilən daha çevik və adaptiv verilənlər integrasiya sistemlərinin tətbiqi tövsiyə olunur. Təhdid cavabını sürətləndirmək üçün mərkəzləşdirilmiş emal infrastrukturuna bağlılığı azaltmaqla kənar hesablama (edge computing) kimi texnologiyalar da nəzərdən keçirilə bilər.

İstifadə edilmiş ədəbiyyat siyahısı

[1] R.H.Aliguliyev, Y.N.Imamverdiyev, and F.J.Abdullayeva, "Artificial intelligence and cybersecurity: New challenges and opportunities" Problems of Information Technology, vol. 12, no. 1, pp. 3–12, 2021.

[2] A.E.Köker, "Siber güvenliğinin stratejik bakış çerçevesinde konumlandırılması," Beykent Üniversitesi Lisansüstü Eğitim Enstitüsü Uluslararası İlişkiler Ana Bilim Dalı, Doktora Tezi, pp. 1–379, 2021.

[3] R.M.Aliguliyev and G.R.Hajirahimova, "Big data analytics in cybersecurity: Trends and challenges in Azerbaijan," Journal of Information Security, vol. 13, no. 2, pp. 45–58, 2023.

[4] H.Chen, Y.Zhang, and D.Wu, "Big data analytics for cybersecurity: A review of trends and challenges," Journal of Information Security and Applications, vol. 70, pp. 103–117, October 2022

[5] Y.Liu, Y.Zhou, S.Hu, and X.Liu, "Explainable AI for cybersecurity: A survey on transparency and interpretability of machine learning models," ACM Computing Surveys, vol. 55, no. 7, pp. 1–38, December 2022

[6] M.A.Al-Garadi, A.Mohamed, A.K.Al-Ali, X.Du, I. Ali, and M.Guizani, "A survey of machine and deep learning methods for Internet of Things (IoT) security," IEEE Communications Surveys & Tutorials, vol. 22, no. 3, pp. 1646–1685, Third Quarter 2020

SIBER GÜVENLİĞİN YAPAY ZEKA ALGORİTMALARI VE BÜYÜK VERİ ANALİTİĞİ İLE GÜÇLENDİRİLMESİ: ZORLUKLAR VE ÇÖZÜMLER

E.A.Muradzade

Mingəçevir Devlet Üniversitesi

Özet: Dijital dönüşüm hız kazandıkça, siber güvenlik giderek karmaşıklaşan tehditlerle karşı karşıya kalmaktadır. Bu çalışma, siber güvenliği güçlendirmek amacıyla yapay zeka algoritmaları ile büyük veri analitiğinin entegrasyonunu ve özellikle şeffaflık sorunlarını incelemektedir. Betimleyici-nitel bir metodoloji kullanılarak literatür, örnek olaylar ve ikincil veriler analiz edilmiş; bu yöntemlerle etkinliğin değerlendirilmesi yapılmıştır. Genel olarak, derin öğrenme teknikleri anomali tespitinde %93 doğruluk oranı sağlamakta ve kural tabanlı yöntemlere kıyasla %18 daha yüksek performans göstermektedir. Örneğin, Spark 500 GB veriyi 35 saniyede işlerken, Hadoop aynı işlemi 60 saniyede tamamlamaktadır. Düzenli güncellemeler, doğruluğu %15 oranında artırmaktadır. Ancak yapay zekanın şeffaflığı ve entegrasyon zorlukları, gerçek zamanlı tehdit tespitini karmaşık hale getirmektedir. Bu çalışma, sistem dayanıklılığını artırmak için stratejiler önererek mevcut alana katkı sağlamaktadır. Önerilen stratejiler arasında şeffaflığı artırmak için açıklanabilir yapay zeka kullanımı ve gelişmiş veri entegrasyon teknikleri yer almaktadır. Gelecekteki araştırmalar, sektöre özgü uygulamalara ve tepki mekanizmalarına odaklanmalıdır.

Anahtar Kelimeler: Siber güvenlik optimizasyonu, yapay zeka (YZ), büyük veri analitiği, tehdit tespiti

УКРЕПЛЕНИЕ КИБЕРБЕЗОПАСНОСТИ С ПОМОЩЬЮ АЛГОРИТМОВ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА И АНАЛИТИКИ БОЛЬШИХ ДАННЫХ: ВЫЗОВЫ И РЕШЕНИЯ

Э.А.Мурадзаде

Мингечаурский государственный университет

Резюме: По мере ускорения цифровой трансформации кибербезопасность сталкивается с всё более сложными угрозами. Данное исследование посвящено изучению интеграции алгоритмов искусственного интеллекта (ИИ) и аналитики больших данных для укрепления кибербезопасности, с особым акцентом на проблемах прозрачности. Используя описательно-качественную методологию, были проанализированы литература, примеры из практики и вторичные данные, что позволило оценить эффективность подходов. В целом, методы глубокого обучения обеспечивают точность обнаружения аномалий на уровне 93%, что на 18% превосходит производительность методов, основанных на правилах. Например, платформа Spark обрабатывает 500 ГБ данных за 35 секунд, тогда как Hadoop выполняет ту же задачу за 60 секунд. Регулярные обновления повышают точность на 15%. Однако проблемы прозрачности ИИ и сложности интеграции затрудняют обнаружение угроз в реальном времени. Настоящее исследование вносит вклад в данную область, предлагая стратегии повышения устойчивости систем, включая использование объяснимого ИИ для обеспечения прозрачности и передовые методы интеграции данных. Будущие исследования должны сосредоточиться на отраслевых приложениях и механизмах реагирования.

Ключевые слова: оптимизация кибербезопасности, искусственный интеллект (ИИ), аналитика больших данных, обнаружение угроз

Elmi redaktor: tex.f.d., dos. A.Mustafayeva

Çapa təqdim edən redaktor: tex.f.d., dos. A.Əliyeva

Daxil olub: 14.04.2025

Çapa qəbul edilib: 20.04.2025