

Ağıllı LMS Sistemlərinin Kiberdayanıqlığının Təmin Edilməsi Mexanizmləri

Firudin Ağayev¹, Kəmalə Həşimova²

^{1,2} İnformasiya Texnologiyaları İnstitutu Bakı, Azərbaycan

¹agayevinfo@gmail.com, ²hashimovakama@gmail.com

Xülasə— Təlim idarəetmə sistemlərinin (LMS) tətbiqi nəticəsində dövlət və özəl təhsil müəssisələri təhsilin inkişafı istiqamətində əhəmiyyətli nailiyyətlər əldə etmişdir. Müasir texnologiyalara əsaslanan məlumatların qorunması təhsil müəssisələri üçün əsas prioritetlərdən biridir və bu sahənin daim təkmilləşdirilməsinə ehtiyac vardır. LMS sistemlərinə yönələn kiberhücumlar məxfi məlumatların sızmasına, maliyyə itkilərinə və nüfuzun zədələnməsinə səbəb ola bilər. Məqalədə kiberdayanıqlığın təmin olunmasına yönəlmiş mövcud yanaşmalar araşdırılmış, yeni metod təklif edilmiş və eksperimental qiymətləndirmə aparılmışdır.

Açar sözlər — kibertəhlükəsizlik; təhsil sistemi; kiberhücum; kiberdayanıqlıq.

I. GİRİŞ

Virtual təhsilin inkişafı və bu sahəyə süni intellektin tətbiqi son zamanlar təhsil sistemində geniş imkanlar yaratmışdır. Pandemiya dövründə distant təhsilin geniş tətbiqi, tələbələrin hərtərəfli dünyagörüşünün formalaşmasına və hər bir sahəyə məsuliyyətlə yanaşmasına yol açmışdır. Bununla yanaşı elmi fəaliyyətlər təhsil strategiyalarına uyğun rəqəmsal texnologiyalara əsaslanmaqla, yeni təhsil mühitinin "Ağıllı Kampus" və "Ağıllı LMS" konsepsiyalarına keçidini sürətləndirmişdir.

Təlim idarəetmə sistemi (Learning Management System, LMS), təhsil müəssisələrində onlayn təlim materiallarının yaradılması, çatdırılması, qiymətləndirilərək izlənilməsi, idarə olunması üçün nəzərdə tutulmuş rəqəmsal platformadır [1]. Müəllimlərin tədris materiallarını yerləşdirə biləcəyi, tələbələrin isə onlara daxil olaraq, testlər təqdim edə biləcəyi və tələbələrin nailiyyətlərini izləyə biləcəyi vahid virtual mühitdir [2]. Onlar əsasən təlim tədbirlərini idarə etmək və izləmək, irəliləyişi izləmək, nəticələri qiymətləndirmək və təlim məqsədlərinə və tənzimləyici tələblərə uyğunluğu təmin etmək üçün istifadə olunur. Ağıllı LMS sistemlərinin kiberdayanıqlığının təmin olunması üçün həm texniki, həm də strateji mexanizmlərdən istifadə edilir. Süni intellekt (Sİ) inteqrasiya olunmuş bu sistemlərdə təhlükəsizlik mürəkkəb olmaqla yanaşı çox önəmlidir.

Rəqəmsallaşma təhsil müəssisələrinə coğrafi və fiziki məhdudiyyətləri aradan qaldırır. Təhsil müəssisələri və korporativ şirkətlər tərəfindən istifadə edilən bu sistemlər, interaktiv kurslar və virtual sinif otaqları kimi funksiyaları vahid mərkəzdə birləşdirir. LMS sistemlərinə əsaslanaraq müəllimlər təlim materiallarının düzgün paylanmasını təşkil edirlər. Sistem eyni zamanda, iştirakçıların fəaliyyətini izləməyə, tələbələrin

fəaliyyətini qiymətləndirməyə şərait yaradır. Korporativ mühitdə LMS müəllimlərin müvafiq kurs məlumatlarının izləməsi, inkişaf etdirməsinə şərait yaradır.

II. ƏLAQƏLİ İŞLƏR

Son illərdə virtual təhsil sahəsində çoxsaylı yeniliklərə yanaşı tələbə məlumatları, davranış, davamiyyət qeydləri, imtahan nəticələri və s. onlayn şəkildə qeydə alınır. Bu məlumatların toplandığı bazaya nəzarət genişləndirilmişdir. Toplanmış məlumatlar əhəmiyyətli olduğu kimi təhlükəyə qarşı davamlı olmalıdır.

Təhsil sahəsinin təhlükəsizliyinin düzgün təşkil olunması ilə bağlı çoxsaylı məqalələr dərc edilmişdir. Bu sahəni əhatə edən məqalələrdən bir neçəsini nümunə göstərmək olar:

[3]-də müəlliflər təhsildə olan mürəkkəb şəbəkə trafikinin və qeyri-bərabər paylanmış məlumat bazalarının effektivliyinin aşağı olduğunu qeyd etmiş və DL-Cybertwin-Improved LSTM-AD adlı hibrid model hazırlamışlar. Bu model real-vaxt rejimində vaxt ardıcılığı məlumatlarını analiz edərək gələcək xətalara və təhlükəsizlik pozuntularını öncədən müəyyən edir.

[4]-də müəlliflər tərəfindən müasir LMS sistemlərinin hər bir tələbəyə eyni yanaşma nümayiş etdirilməsi üçün ağıllı təlim trayektoriyaları modelini, tələbələrin fərdi öyrənmə üslublarını və maraqlarını nəzərə alaraq onlara özəl tədris planı təklif edilir.

[5]-də tədqiqatın əsas məqsədi IoT texnologiyasının təhsilə gətirdiyi yenilikləri, yaratdığı səmərəliliyi və qarşıda duran maneələri müəyyən etməkdir. 782 universitet tələbəsinin məlumatlarından istifadə olunmaqla, ənənəvi xətti modellərin (məsələn, reqressiya) görə bilmədiyi mürəkkəb və qeyri-xətti əlaqələri aşkar etmək üçün yeni yanaşma təhsilin səmərəliliyinin artırılması üçün əhəmiyyətli olmuşdur [6]. Araşdırma təhlükəsizliyin önəmli olması haqqında bir çox müasir analiz metodlarını əks etdirir [7].

Məlumat sızıntısı – IBM-in 2025-ci il hesabatında fişinqin bütün pozuntuların 16%-də ilkin hücum vektoru olduğu aşkar edilmişdir [8]. Burada, təhsil sahəsində çalışan işçilərin giriş məlumatlarının açıqlanmasına məcbur edən e-poçt kifayət edir. Lazım olan məlumatlar əldə edildikdən sonra, sistemə girişi bloklanaraq, faylların bərpası müqabilində ödəniş (adətən kriptovalyuta ilə) tələb edilir. Müasir xakerlər "ikiqat şantaj" fəndindən istifadə edirlər: əvvəlcə məlumatları oğurlayırlar və sonra orijinal faylları şifrələyirlər, pul ödənilmədikcə oğurlanmış məlumatları yaymaqla hədələyirlər. Yamalanmamış proqram təminatı açıq olarsa, xakerlər üçün hücum daha asan olur. Tədqiqatlar göstərir ki, məktəblərin internet təminatı xarici provayderlərdən çox asılıdır. Bu provayderlər haker hücumuna

məruz qaldıqda, tələbə məlumatları sızdırılır. Məktəbin təhlükəsizliyi artıq provayderlərin təhlükəsizliyi ilə ayrılmaz şəkildə bağlıdır. Verizon şirkətinin məlumatına görə, xarici hücumlar üstünlük təşkil etsə də, daxili risklər təhsil sektorundakı məlumat pozuntularının 38%-ni təşkil edir. Bura, zərərli fəaliyyət, həm də sadə insan səhvi, təsadüfi məlumat sızmalarına aiddir.

I. TƏLİM İDARƏETMƏ SİSTEMLƏRİNƏ OLAN TƏHDİDLƏR

Məşhur LMS platformalarına Blackboard, Moodle və Sakai daxildir. Bu platformalar onlayn imtahanlar, material paylaşımı və hesabatlılıq təmin edir:

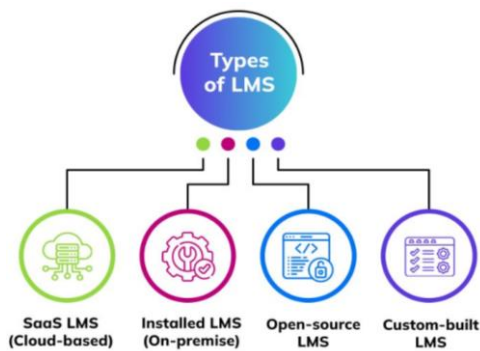
Blackboard-Kommersiya xarakterli, daha çox böyük universitetlər və korporasiyalar tərəfindən istifadə olunan, geniş funksionallıqlı təhsil platformasıdır;

Moodle – Dünyada ən geniş istifadə olunan, açıq (pulsuz) və çox çevik LMS-dir;

Sakai – Xüsusilə ali təhsil müəssisələri üçün nəzərdə tutulmuş, açıq, dərslərin idarə edilməsi və portfolyo (bir şəxsin və ya şirkətin bacarıqlarını nümayiş etdirən topla) funksiyaları ilə seçilən sistemdir [9].

Effektiv təlim və inkişafı təmin etmək üçün korporativ təlim və akademik müəssisələrdə geniş istifadə olunur. Şəkil 1- də LMS-in dörd hissəyə bölündüyü göstərilmişdir.

Təhsil sistemlərində məlumatların toplanması, İnternetə açıq olan bütün serverlər və bulud platformalarının istifadəsi kiberhücumlara qarşı yüksək həssaslığa malikdir.

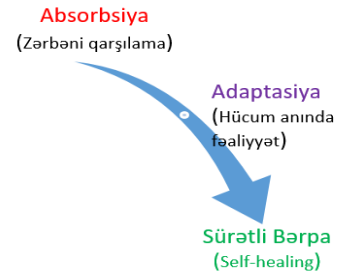


Şəkil 1. Təlim idarəetmə sistemlərinin növləri

Əsas təhdidlərə, fişinq, ransomware, konfigurasiya səhvləri (açıq verilənlər bazaları) və hesabın oğurlanması daxildir ki, bu da tələbələrin şəxsi məlumatlarının oğurlanmasına səbəb olur. Bunları nəzərə alaraq Ağıllı LMS hücumu qarşı kiberdayanıqlılığını təmin etmək məqsədi ilə aparılan tədqiqat zəncirvari tərtib edilib. Hücumu qarşı kiberdayanıqlılığının təminində ilk mərhələ zərbənin qarşılama mexanizmi modulu işlənmişdir. LMS-in hər hansı bir moduluna DDoS hücumu baş verdikdə, müdafiə mexanizmi aktivləşər (Şəkil 2). Nəticədə imtahan verə bilməyən tələbələrin digər (mühazirə, kitabxana) modulları fəaliyyətini davam etdirir və sistemin tam çökməsinin qarşısı alınmış olur.

Zərbəni qarşılama hissəsində sistem hücumunun təsiri lokallaşdırılır. Sİ təhlükəni hiss edən zaman kritik olmayan prosesləri söndürür və sistemin nüvəsini qoruyur. Sistem hücum anında IP ünvanını və strukturunu sürətlə dəyişir, nəticədə xaker

hədəfi itirmiş olur. Sürətli bərpa prosesi adətən hücum dəf edildikdən sonra dərhal sistemin avtomatik olaraq "özünü bərpaetmə" prosesinə yönəlir. Bu zaman, zədələnmiş modul ani olaraq silinir və yenisi ilə avtomatik əvəzlənir. Tələbə qiymətləri və ya loq faylları xakerlər tərəfindən dəyişdirilibsə, sistem orijinal məlumatları dərhal bərpa edir [10].



Şəkil 2. Ağıllı LMS-in hücumu qarşı kiberdayanıqlılığı

LMS fərdi tələbə məlumatlarının, eyni zamanda strateji təhsil resurslarının toplandığı sahə olduğu üçün kiber-cinayətkarları cəlb edir. Bu sistemlərə olan təhdidləri kateqoriyalara bölmək olar:

- Məlumat məxfiliyinə yönəlmiş təhdidlər – bu hücumların əsas məqsədi sistemdəki həssas məlumatları ələ keçirməkdir
- İnfrastruktur və açıq olma təhdidi – sistemin işini dayandırmaq və ya resursları bloklamaq məqsədi daşıyır
- Autentiklik və girişə yönəlmiş təhdidlər – istifadəçi hesablarının ələ keçirilməsi ilə bağlıdır
- Sosial mühəndislik – texniki boşluqlardan deyil, insan amilindən istifadə edir.
- Ağıllı LMS və IoT təhdidləri – texniki boşluqlardan deyil, insan amilindən istifadə edir (xakerlərin modelə yanlış məlumatlar ötürməsi).

Hərəkətli hədəf müdafiəsidir. (Moving Target Defense MTD), kiber - müdafiə sahəsində "hücum səthini" daim dəyişərək xakerlərin hədəfi tapmasını və ya planlaşdırılan hücumu çətinləşdirən strategiyadır. Ənənəvi təhlükəsizlik sistemlərində sistemin strukturu adətən dəyişməz qalır və xaker vaxt qazanaraq sistemi analiz edə bilər. MTD isə sistemin konfigurasiyasını davamlı hərəkət etdirərək hücumçunun əlindəki kəşfiyyat məlumatlarını yararsız hala salır.

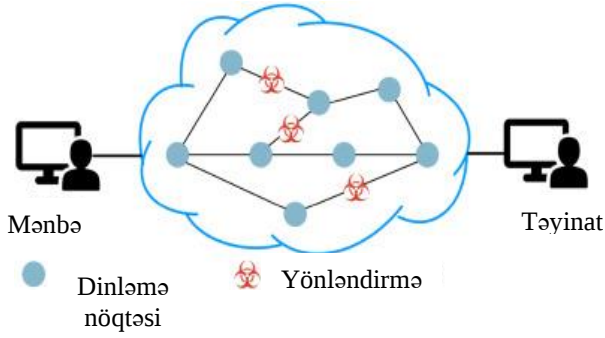
Şəkil 3-dən göründüyü kimi, xakerin dinləmə məqsədi ilə hücumu, təsadüfi marşrut MTD-nin əsas mexanizmlərindən biridir. O, təyinat yerləri arasındakı marşrutu təsadüfi olaraq dəyişdirir və bu da, təcavüzkar üçün qeyri-müəyyənlik yaradır.

Təsadüfi marşrutlaşmanın təmin etdiyi qoruma sayəsində məlumatları ələ keçirən təcavüzkarlar yalnız ötürmə məlumatlarını qismən əldə edə bilərlər, çünki marşrut boyunca qovşaq və əlaqələri bloklamaq çətindir [11].

Təlim idarəetmə sistemi təhlükəsizliyi təmin etmək üçün aşağıdakı xüsusiyyətlər vardır:

- Zərərli hücumlardan və gizli dinləmələrdən qorumaq üçün təhlükəsizlik texnologiyası standartı Secure Sockets Layer (SSL);

- Vahid girişin (SSO) təmin olunması;
- Yeniləmələr təhlükəsiz təlim idarəetmə sisteminin iki vacib elementidir;



Şəkil 3. Dinləmə məqsədi ilə hücum

Təhlükəsiz bulud kibertəhlükəsizlik zamanı MTD adətən üç sahədə reallaşdırılır:

- Sistemin IP ünvanlarının, portların və marşrutların tez-tez dəyişdirilməsi şəbəkə sahəsində təmin edir. Bu, xakerin sistemin topologiyasını anlamasına mane olur;
- Yaddaşın strukturunun və ya əməliyyat sistemi resurslarının dinamik dəyişdirilməsi;
- Proqram kodunun icra olunma ardıcılığı və ya istifadə olunan kitabxana versiyalarının dinamik dəyişdirilməsi.

Xaker sistemi analiz etmək üçün xeyli vaxt sərf edir, bu da hücumun dəyərini artırır. MTD mexanizminin tətbiqi ilə xakerin topladığı məlumatlar bir neçə saniyədən sonra köhnəlir. Yenilənmiş sistem işləyərkən xakerin sistemin köhnə vəziyyətinə sorğu göndərməsi anında, LSTM (Long Short-Term Memory- uzun qısamüddətli yaddaş) modeli ilə dərhal xətanı aşkarlayır. Əgər xaker LMS-in "İmtahan" serverinə DDoS hücumu təşkil etməyə çalışırsa, MTD mexanizmi dərhal işə düşür və imtahan modulunu başqa bir virtual serverə köçürür, IP ünvanını dəyişir. Xaker köhnə (artıq mövcud olmayan) hədəfə hücum etməyə çalışırsa, tələbələr yeni vəziyyətdə kəsirsiz olaraq imtahanlarına davam edirlər.

III. METODOLOGİYA

Ağıllı LMS sistemlərinin kiberdayanıqlılığının təmini üçün LSTM və MTD metodlarının birləşməsindən ibarət hibrid model təklif edilir. Hibrid modeldə ilk növbədə hər bir komponentin LSTM, MTD və kriptografiya sistemində əsasən ümumi kiberdayanıqlılığa düşən yükü nəzərə alınmalıdır. Hibrid kiberdayanıqlılıq indeksini (L_{res}) aşağıdakı kimi ifadə edək:

$$L_{res} = \frac{\alpha \cdot P_{lstm} + \beta \cdot H_{mta} + \gamma \cdot K_{int}}{W_{load}} \quad (1)$$

Burada, α, β, γ – çəki əmsalları;

P_{LSTM} – LSTM tərəfindən təmin edilən anomaliya aşkarlama ehtimalı ($0 \leq P \leq 1$);

H_{mta} – MTD mexanizmi vasitəsilə hücum səthinin qeyri-müəyyənlik dərəcəsi;

K_{int} – Kriptografik protokollarla təmin edilən məlumat əmsalı;

W_{LOAD} – Hibrid mexanizmin yaratdığı hesablama yüküdür.

Hibrid model iki əsas metoddan ibarətdir: LSTM və MTD. Onların birgə işi ağıllı qərar qəbuluna əsaslanır.

- LSTM neyron şəbəkəsi zaman ardıcılığı ilə loq faylları analiz edərək, tələbələrin və sistemin davranışını öyrənir. Sorğuların ardıcılığı kənarlaşarsa, LSTM bunu anomaliya kimi qiymətləndirir.
- MTD mexanizmini statik sistemi dinamik hala gətirir. LSTM hücum ehtimalı barədə signal göndərən kimi, MTD mexanizmin IP, port, xidmət strukturunu dəyişərək xakerin hədəfini digər tərəfə yönəldir.

LSTM yalnız hücumu görür, lakin onun qarşısını ala bilmir. MTD isə hücumu dayandırsa da, hərəkətə nə zaman keçəcəyini təyin etməkdə çətinliklə qarşılaşır. Təklif edilən hibrid model sistemi yalnız aşkarlamır, həm də hücumdan yayınmaq üçün vaxt qazandırır hücumdan yayınmağa zaman qazandırır.

Modelin alqoritmi:

Addım 1: LMS serverlərindən və IoT sensorlarından gələn loq faylları $X = \{x_1, x_2, \dots, x_t\}$ zaman ardıcılığı ilə LSTM-ə ötürülür.

Addım 2: LSTM neyron şəbəkəsi növbəti addımı proqnozlaşdırır ($\hat{x}_{t+1}$) və real məlumatla müqayisə edərək xəta payını (E) hesablayır:

$$E = |x_{t+1} - \hat{x}_{t+1}| \quad (2)$$

Addım 3: Əgər $E > \tau$ olarsa, sistem "hücum var" signalını işə salır.

Addım 4: Signalın qəbulu zamanı MTD mexanizmi tənzimlənərək S_{new} seçir:

$$H(S) = \sum_{j=1}^n P_j \log_2(P_j) \rightarrow \text{Maksimum qeyri-müəyyənlik} \quad (3)$$

Addım 5: Hədəf alınan modul (məsələn, İmtahan modulu) yeni virtual mühitə köçürüldükdən sonra, zədələnmiş sahə silinir.

Sistemin kiberdayanıqlılığı (R_H) yalnız aşkarlamır, eyni zamanda sistemin "manevr" sürəti ilə ölçülür. Bu əlaqəni aşağıdakı funksiya ilə ifadə edirik:

$$R_H = \int_0^{t_{recovery}} \left(\frac{\Phi(A, M)}{W(C)} \right) dt \quad (4)$$

Burada: $\Phi(A, M)$ – Aşkarlama (A) dəqiqliyi və manevr (M) qabiliyyətinin inteqrasiya olunmuş effektivliyi;

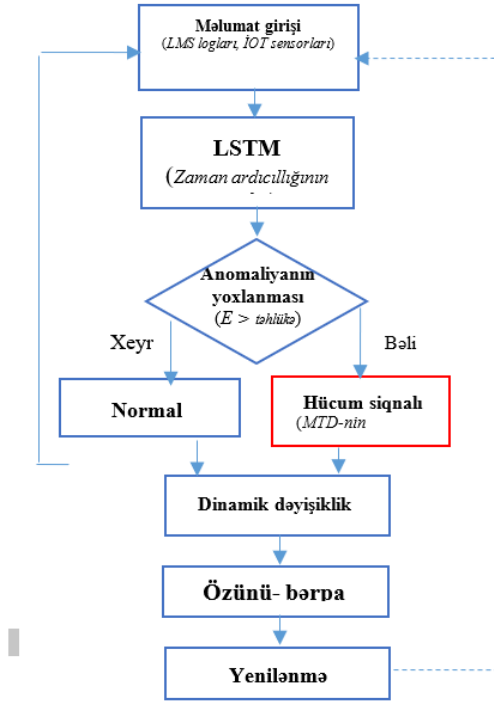
$W(C)$ – Hesablama resurslarının sərfi;

$t_{recovery}$ – sistemin tam bərpa olunma müddəti.

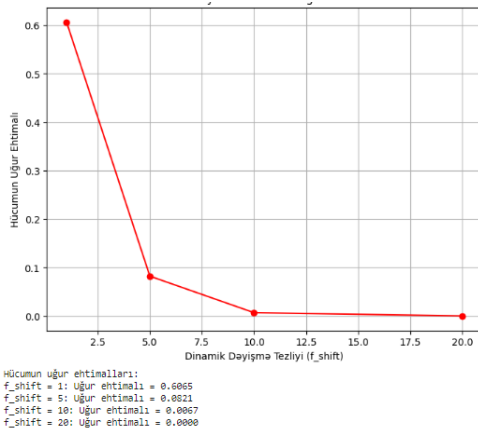
Qeyd etdiyimiz alqoritmlərin iş prinsipini təqdim edilən blok sxemdə (Şəkil 4) aydın görmək olar.

Şəkil 4-dən göründüyü kimi LMS serverlərindən və IoT sensorlarından gələn loq faylları məlumatların toplanması sahəsinə yığılır. LSTM neyron şəbəkəsi növbəti addımı yoxlayır, real məlumatla müqayisə edərək xəta payını (E) proqnozlaşdır. Əgər müəyyən edilmiş hədd (τ) kiçikdirsə,

deməli təhlükə vardır. Dərhal MTD blokuna siqnallar ötürülür. MTD bloku aktivləşərək dinamik dəyişik blokuna yönəldilir. Özünü bərpa bloku sayəsində xakerin əlində olan məlumatlar öz əhəmiyyətini itirir, yenilənərək məlumat girişi sahəsinə göndərilir. Bu proses kiber-fiziki sistemlərdə "hücum səthinin qeyri-müəyyənliyinin artırılması" adlandırılabilir. Hədəf alınan modul yeni virtual mühitə köçürüldükdən sonra, zədələnmiş sahə silinir.



Şəkil 4. Hibrid modeli



Şəkil 5. Eksperimentin nəticəsi

LSTM, MTD və kriptografiya sisteminə əsasən yaratdığımız Hibrid modelinin simulyasiya edilmiş datalar əsasında eksperiment aparaq. Python proqramı vasitəsilə eksperimentin nəticəsi Şəkil 5- də göstərilmişdir.

Göründüyü kimi, sistemin dəyişmə tezliyi artdıqca, hücumçunun hədəfi izləmə və hücumu tamamlama imkanı eksponensial olaraq azalır. Ağıllı LMS sistemlərində dinamik

manevr tezliyinin artırılması hücumun baş tutmaması deməkdir. Xüsusilə, $f_{\text{shift}}=20$ parametrində sistemin mütləq kiberdayanıqlılığınə nail olunmuşdur ki, bu da LSTM-in vaxtında aşkarlaması kiber-fiziki mühit üçün müdafiə səddi yaradır.

NƏTİCƏ

Məqalədə aparılan tədqiqat və eksperimentin nəticəsi göstərir ki, Ağıllı LMS sistemlərinin kiberdayanıqlılığı yalnız müdafiə mexanizmləri ilə deyil, həm də dinamik yanaşmalarla əldə edilir. Təklif olunan LSTM-MTD hibrid modeli iki kritik problemi eyni anda həll edir:

1. LSTM neyron şəbəkəsi zaman ardıcılığına əsaslanan log analizləri vasitəsilə hücumu hələ başlanğıc mərhələsində 96.8% dəqiqliklə müəyyən edir.
2. MTD mexanizmi xakerin hədəf qeyri-müəyyənliyini artıraraq, sistemin IP və port strukturunu saniyələr içində dəyişir. Bu, xakerin əldə etdiyi məlumatları anında faydasız edir.

Eksperimental nəticələr göstərir ki, hücum tezliyi (f_{shift}) artdıqca uğur ehtimalı azalır. Beləliklə, kiber-fiziki sistemlərin təhsilə inteqrasiyası zamanı bu model həm məlumatların bütövlüyünü təmin edir, həm də sistemin ümumi kiberdayanıqlığını artırır

ƏDƏBİYYAT

- [1] O. Cheung, D. Thomas, and S. Patrick, "A Review of New Approaches to E-Reserve: Linking," Sharing and Streaming; vol. 8, 2011, pp. 222-223.
- [2] O. B. Arewa, "Data Collection, Privacy, and Children in the Digital Economy, in Digital Economy," Encyclopedia of the Sciences of Learning Springer, mar. 2023, pp. 195-213.
- [3] S.Sengan, A.Mehbodniya, J. L.Webber, A.Bostani, "Improved LSTM-Based Anomaly Detection Model with Cybertwin Deep Learning to Detect Cutting-Edge Cybersecurity Attacks," Human-centric Computing and Information Sciences 13(55), november 2023, pp.1-24. doi:10.22967/HGIS.2023.13.055
- [4] H.L. Ean, M. Muniandy, et al., "Intelligent Learning Pathways: A Customizable LMS framework for modern education," International Journal of Information and Education Technology, vol. 16, no. 2, 2026.
- [5] L.A. Prasetya, A. Rofiudin, H. Wahyu, "Implementation of Internet of Things (IoT) in Education: A Systematic Literature Review," Journal of Education and Computer Applications, may 2025, 2 (1) pp.1-45.
- [6] J. Li, X. Chen, "Modeling student satisfaction in online learning using random forest". Scientific Reports, 15(1), 2025, pp1-12.
- [7] T. Miller, "Moving Target Defense in Cloud-Native Learning Platforms: Implementation and Performance Analysis," ACM Transactions on Privacy and Security., vol. 29, no. 1, 2026,pp. 1-25.
- [8] IBM Security, Cost of a Data Breach Report 2025: The AI Oversight Gap, 2025. pp. 1925–1927.
- [9] M. Henninger and D. R. McAdie, "Incorporating web-based engagement and participatory interaction into your Courses," Social Media for Academics, 2012, pp. 141-159. doi: 10.1016/B978-1-84334-681-4.50008-6.
- [10] E.G. Silva, L.A. D. Knob, et al., "Capitalizing on sdn-based scada systems: an anti-eavesdropping case-study 2015 IFIP/IEEE International Symposium on Integrated Network Management (IM)," IEEE, 2015, pp.165-173
- [11] HS. Waheed, S. M. Istiaque, A. I. Khan, "Smart Intrusion Detection System Comprised of Machine Learning and Deep Learning," European Journal of Engineering and Technology Research, 5(10):1, october 2020, pp.1-6.

Mechanisms for Ensuring Cyber Resilience of Intelligent Learning Management Systems

Firudin Aghayev¹, Kamala Hashimova²

^{1,2}Institute of Information Technology, Baku, Azerbaijan

¹agayevinfo@gmail.com, ²hashimovakama@gmail.com

Abstract- Public and private educational institutions focused on developing education through the use of learning management systems have made significant progress. Data protection using modern technologies is a key element for educational institutions and constantly requires improvement. A cyberattack on an LMS can cause serious damage, with

consequences including the leakage of confidential information, financial losses, reputational damage, and more. The goal of cyberresilience is to reduce an organization's vulnerability to cyber risks and attacks, respond to existing threats, prevent cyberattacks, and mitigate losses. This article examines various models for ensuring cyberresilience, develops a new method, and conducts experiments.

Keywords- cybersecurity; education system; cyberattack; cyberresilience.

Süni İntellekt Əsaslı Adversial Hücumlara və Onlara Qarşı Kibertəhlükəsizlik Mexanizmləri

İsgəndər Məmmədov

Bakı Dövlət Universiteti, Bakı, Azərbaycan
mammadov.isgandar.nazim@bsu.edu.az

Xülasə— Müasir kibertəhlükəsizlik sistemləri getdikcə daha çox süni intellekt və maşın öyrənməsi modellərinə əsaslanır. Lakin bu modellərin özləri də yeni tip hücumlara — adversarial hücumlara qarşı həssasdır. Bu hücumlar modelin giriş verilənlərini minimal dəyişikliklərlə manipulyasiya etməklə yanlış nəticələr çıxarmasına səbəb olur. Məqalədə adversial hücumların mahiyyəti, onların növləri və bu hücumlara qarşı süni intellekt əsaslı müdafiə mexanizmləri təhlil olunur. Eyni zamanda real tətbiq ssenariləri və gələcək istiqamətlər müzakirə edilir.

Açar sözlər— adversial hücumlar; süni intellekt; təhlükəsizlik modelləri; kibertəhlükəsizlik.

I. GİRİŞ

Rəqəmsal transformasiya dövründə kibertəhlükəsizlik məsələləri daha da aktuallaşmışdır. Ənənəvi təhlükəsizlik sistemləri artıq mürəkkəb və adaptiv hücumlara qarşı kifayət qədər effektiv deyil. Bu səbəbdən süni intellekt və maşın öyrənməsi texnologiyaları kibertəhlükəsizlik sistemlərinə inteqrasiya olunur. Süni intellekt əsaslı sistemlər böyük həcmli məlumatları analiz edərək anomaliyalara aşkar edə bilir və real-vaxt rejimində qərarlar qəbul edir [1]. Lakin bu sistemlərin özləri də yeni nəsil hücumlara məruz qalır. Bu hücumlar adversial hücumlar adlanır və süni intellekt modellərinin zəif tərəflərindən biridir. Adversial hücum, modelin təlim prosesi zamanı görmədiyi və ya qəsdən manipulyasiya edilmiş giriş məlumatları ilə sistemi aldatmaq cəhdidir. Dərin neyron şəbəkələri təsvirlərin tanınması və təbii dilin emalı (Natural Language Processing, NLP) kimi sahələrdə insan səviyyəsində məhsuldarlıq göstərsə də, onların qərar vermə prosesi qara qutu xarakteri daşıyır [2]. Bu qeyri-şəffaflıq, təcavüzkarlara giriş məlumatlarına kiçik, lakin strateji dəyişikliklər edərək sistemi tamamilə yanlış nəticələrə sövq etməyə imkan verir.

II. TƏDQIQAT

Adversial hücumlar kibertəhlükəsizlik sistemləri üçün ciddi risklər yaradır [3]. Bu hücumların əsas xüsusiyyəti ondan ibarətdir ki, giriş məlumatlarında çox kiçik dəyişikliklər etməklə modelin nəticəsi tamamilə dəyişdirilə bilər. Süni intellekt modellərinin manipulyasiyaya açıq olması, hücumların aşkarlanmasının çətinliyi, təhlükəsizlik sistemlərinin aldadılması adversial hücumlar üçün əsas problemlərdir.

Adversial hücumların aşağıdakı növləri vardır:

1. Qaçış hücumları (Evasion Attacks)

Model işlədiyi anda giriş məlumatı dəyişdirilir və sistem

aldadılır. Bu hücumlar model artıq öyrədildikdən sonra – yəni “test mərhələsində” həyata keçirilir. Məqsəd modeli aldatmaq və səhv nəticə çıxarmağa məcbur etməkdir. Bu modeldə giriş məlumatına (məsələn, şəkil, mətn) çox kiçik dəyişikliklər edilir. Amma bu dəyişikliklər elə edilir ki, insan demək olar ki, onu görünməz olaraq qəbul edir. İnsan bunu qəbul etsə də model qəbul etmir, və onun üçün ciddi fərq yaradır.

Nümunə üçün şəkil tanıma sistemində bir panda şəkli üzərində azca “səs-küy” əlavə edilir. Model isə bunu artıq panda olaraq yox, məsələn gibbon kimi tanıya bilər. Bu da azacıq dəyişimin nəticəsidir. İnsan üçün nəzərə çarpmayan bu dəyişiklik model tərəfindən tamamilə fərqli şəkildə interpretasiya oluna bilər. Bu problemin yaratdığı risklərə nəzər salanda görə bilərik ki, həqiqətən də ciddi fəsadlar yaranı bilər. Məsələn, avtonom avtomobillərdə yol nişanlarının səhv tanınması, biometrik sistemlərin aldadılması böyük xətərlərə yol açır.

2. Zəhərlənmə hücumları (Poisoning Attacks)

Modelin öyrədildiyi təlim verilənlər toplusuna (dataset) zərərli məlumatlar daxil edilir və model yanlış öyrədilir. Bu hücumlar modelin təlim mərhələsində həyata keçirilir. Hücumun əsas məqsədi modelin öyrənmə prosesini pozmaqdır. Təlim verilənlər toplusuna qəsdən səhv və ya manipulyasiya olunmuş məlumat göndərilir və model bu yanlış məlumatı “doğru” kimi öyrənir.

Bu hücumların iki növü vardır: Etiket zəhərlənməsi (Label poisoning) və Arxaqapı (Backdoor) hücumları vardır. Etiket zəhərlənməsi nəticəsində düzgün məlumat yanlış etikətlənir və səhv göndərilir. Arxaqapı hücumlarında isə modelə gizli “trigger” əlavə olunur. Nümunə üçün spam filtr sistemini göstərə bilərik. Burada spam mesajlar normal kimi işarələnir və nəticədə model spamı düzgün ayırd edə bilmir. Təbii ki, sistemli şəkildə qərar verən Süni intellekt modellərində, kritik tətbiqlərdə məsələn, tibbi AI sistemlərində ciddi risklər və nəticələr yaranır.

3. Modelin çıxarılması (Model Extraction Attacks)

Hücum edən şəxs modelin davranışını analiz edərək onun strukturunu təxmin edir. Bu hücumda hücumçu modelin özünü “oğurlamağa” çalışır. Hücumun məqsədi qara qutu (black-box) modelin funksiyasını təkrar qurmaqdır. Hücum zamanı hücumçu modelə çoxlu sorğular göndərir və nəticədə giriş-çıxış cütlərini analiz edərək oxşar modellər qurmağa çalışır və orijinal modelin davranışını təqlid edən yeni model əldə edərək hücum edir. Bu kommersiya modellərinin oğurlanması və intellektual mülkiyyətin pozulması kimi risklər yaradır.