



TARMOQ HUJUMLARINI ANIQLASHDA SUN'IY INTELLEKT ALGORITMLARINI QO'LLASH

**Ibrohimov Azizbek
Ravshonbek o'g'li**

“Kiberxavfsizlik markazi” DUK bo'lim boshlig'i

**Haydarov Elshod
Dilshod o'g'li**

Muhammad al-Xorazmiy nomidagi TATU, kafedra mudiri

Annotatsiya

Ushbu maqolada tarmoq hujumlarini aniqlashda sun'iy intellekt algoritmlarini qo'llash samaradorligi tahlil qilindi. Mashinani o'rganish va chuqur o'rganish modellarining afzalliklari, ularning signaturali va anomaliya asosidagi tizimlarda qo'llanish imkoniyatlari ko'rsatildi. Tadqiqot natijalari gibridd yondashuvlar eng yuqori aniqlikni ta'minlashini ko'rsatdi. Kelgusida graf neyron tarmoqlar va explainable AI texnologiyalaridan foydalanish istiqbollari tavsiya etildi.

Kalit so'zlar:

tarmoq hujumlari, sun'iy intellekt, mashinani o'rganish, chuqur o'rganish, anomaliya, signaturali tizim, IDS, GNN.

Аннотация: В статье проанализирована эффективность применения алгоритмов искусственного интеллекта для обнаружения сетевых атак. Рассмотрены преимущества методов машинного и глубокого обучения, их использование в сигнатурных и аномальных системах обнаружения. Результаты исследования показали, что гибридные подходы обеспечивают наибольшую точность. В перспективе рекомендуется использование графовых нейронных сетей и технологий объяснимого искусственного интеллекта.

Ключевые слова: сетевые атаки, искусственный интеллект, машинное обучение, глубокое обучение, аномалия, сигнатурная система, IDS, GNN.

I. KIRISH

Bugungi global raqamli jamiyatda axborot va kommunikatsiya texnologiyalari hayotimizning deyarli barcha jabhalarida markaziy rol o'ynamoqda. Internet-tarmoqlar, korporativ infratuzilmalar, sanoat nazorati tizimlari, bulutli xizmatlar va IoT qamrovi kengaygan sari, ushbu sistemalarga nisbatan hujumlar soni va murakkabligi ham keskin oshdi. Tarmoq hujumlari — bu tarmoqqa yoki u bilan bog'liq infratuzilmaga ruxsatsiz kirish, ma'lumot o'g'irlash, xizmatni bloklash (masalan DDoS), skanerlash yoki boshqa zararli faoliyatlarni amalga oshirishga qaratilgan xatti-harakatlar to'plamidir. Ular tizimning maxfiyligi, yaxlitligi va mavjudligini buzishi mumkin, xususan muhim infratuzilmalar — moliya, energetika, sog'liqni saqlash sohalaridagi axborot tizimlari uchun xavf tug'diradi.

An'anaviy IDS (intrusion detection system) va IPS (intrusion prevention system) yechimlari, asosan, imzolarga (signature) yoki qoidalarga tayanadi — ya'ni ma'lum hujum tiplariga moslangan shablonlarni tekshiradi. Biroq hujum texnikasi doimiy ravishda rivojlanmoqda, zero-day hujumlari paydo bo'lmoqda, va hujumchilarning shovqinli (obfuscated) usullarni qo'llash qobiliyati ortmoqda. Shuning uchun faqat signaturaga tayanadigan tizimlar kamdan-kam hollarda barcha xavflarni aniqlay oladi.

Shu kontekstda, sun'iy intellekt (AI) algoritmlari — mashinani o'rganish (ML), chuqur o'rganish (DL) va bog'lanishli (graph) tarmoqlar — tarmoq hujumlarini aniqlashda kengroq va puxta tahlil imkoniyatlari bilan diqqat markaziga chiqdi. AI asosida ishlab chiqilgan yondashuvlar trafikdagi murakkab naqshlarni o'rganishi, normal holatdan og'ishlarni aniqlashi va yangi hujumlarni oldindan proaktiv tarzda aniqlashga qodir. Shu bilan birga, ularning operatsion muhitga tatbiq etilishi, real-vaqt ishlashi, noto'g'ri signal (false positive) darajasi hamon muammolar sirasida.

Ushbu maqolaning maqsadi — tarmoq hujumlarini aniqlash kontekstida sun'iy intellekt algoritmlarini qo'llash imkoniyatlarini ilmiy jihatdan ko'rib chiqish, sohadagi metodologik yondashuvlarni tahlil qilish va adabiyotlar sharhi orqali mavjud bo'shliqlarni aniqlashdir. Mazkur muammo dolzarbligi, chunki raqamli infrastruktura kengayishi bilan hujumlarning chastotasi va murakkabligi ortmoqda, korporativ va davlat darajasida xavfsizlik zaifliklari yuritilmoqda, shuningdek operatsion real-vaqt aniqlash tizimlari uchun yuqori talab mavjud.

II. TADQIQOT METODOLOGIYASI

Tadqiqot jarayonida quyidagi tahliliy metodologik yondashuvlar qo'llanadi:

Adabiyotlar tarixi va sistematik tahlil (systematic literature review). Avvalo, soha bo'yicha nashr qilingan maqolalar, sharhlar va tadqiqot ishlarini aniqlash, ularni ma'lum mezonlar bo'yicha saralash va tahlil qilish nazarda tutiladi. Bu bosqichda so'rov formulalari ishlab chiqiladi, tegishli ma'lumot bazalarida (IEEE Xplore, Web of Science, SCOPUS, arXiv) qidiruv amalga oshiriladi, tegishli maqolalar tanlanadi va sifat jihatdan baholanishi (masalan, oborotli jurnal bo'lishi, sitatlar soni, metodologiya sifati) ko'zdan kechiriladi. Masalan, "Survey on intrusion detection systems: techniques, datasets and challenges" maqolasi tarmoq hujumlari aniqlash usullari va datasetlarini tartibga solib beradi.

Tanlangan adabiyotlar ichida ishlatilgan metodlar, algoritmlar, datasetlar, baholash metrikalari, operatsion muhitdagi implementatsiya tajribasi tahlil qilinadi. Bu bosqichda adabiyotlardan olingan ma'lumotlar jamlanadi, umumiy tendensiyalar va muammo nuqtalari aniqlanadi. Shu orqali tadqiqot uchun «qaysi algoritmlar eng ko'p qo'llanilgan», «qaysi datasetlar eng keng tarqalgan», «qaysi baholash metrikalari dominant» degan savollarga javob izlanadi.

O'zaro metodologiyalarni taqqoslash (masalan, supervayzlangan vs unsupervayzlangan usullar), ularning kuchli va zaif tomonlari aniqlanadi. Shuningdek, amaliy tatbiqiy jihatdan modelni real-vaqt tizimga joylash, ma'lumotlarni balanslash (class imbalance), drifti (concept drift), resurs iste'moli va kechikish (latency) masalalari ham ko'rib chiqiladi. Masalan, adabiyotlarda ML modelidagi umumlashuv cheklovlari va real trafikdagi murakkabliklar muammo qilib ko'rsatilgan.

O'rganilgan adabiyotlar tahlilidan so'ng, xozirgi holatda qaysi yo'nalishlar kam o'rganilganligi, qaysi amaliy muammolar mavjudligi aniqlanadi (masalan, zero-day

hujumlarini aniqlash, real-vaqt operatsion muhitda model joylashuvi, explainability, adversarial hujumlarga chidamlilik). Ushbu tahlil asosida tadqiqotning kelgusi yo‘nalishlari va tavsiyalar shakllantiriladi.

Tahlil natijalari tartiblangan holda tayyorlanadi: qaysi yondashuvlar eng keng qo‘llanilgan, ularning muvaffaqiyat darajasi va muammolari. Bu bosqichda jadval, grafiklar, statistik summarizatsiya qo‘llanilishi mumkin, lekin hozirgi maqola ichida asosan tushuntirish shaklida keltiriladi.

Ushbu metodologiya tadqiqotni ilmiy, tizimli va izchil shaklda olib borish imkonini beradi hamda ham mavjud adabiyotlar asosida, ham amaliy tahlillar asosida aniq natijalarga erishishni ta‘minlashga xizmat qiladi.

III. ADABIYOTLAR SHARHI

So‘nggi yillarda tarmoq hujumlarini aniqlash sohasida sun‘iy intellekt texnologiyalaridan foydalanish bo‘yicha juda ko‘p ilmiy tadqiqotlar olib borilmoqda. Dastlabki davrda hujumlarni aniqlash asosan *signaturali* (*imzoli*) tizimlarga asoslangan edi. Bunday tizimlar har bir hujum turi uchun oldindan yozilgan qoida yoki shablonlar orqali ishlaydi. Masalan, Snort yoki Suricata kabi tizimlar shunday ishlaydi. Ammo yangi hujum turlarining tez paydo bo‘lishi, ularning shakli va usullari murakkablashgani sababli bu yondashuvlar yetarli bo‘lmay qoldi. Chunki ular ilgari kuzatilmagan (“zero-day”) hujumlarni aniqlay olmaydi [1].

Shu sababli, ilmiy tadqiqotlar markazi sun‘iy intellekt (AI) algoritmlariga qaratildi. Bu yondashuvlar tarmoqdagi trafikni avtomatik tahlil qilib, normal va anomal holatlarni farqlay oladi. Ular qoidalarga tayanmaydi, balki o‘rganilgan ma‘lumotlar asosida mustaqil qaror chiqaradi [2].

Masalan, Vinayakumar (2019) o‘z ishida chuqur o‘rganish (Deep Learning) modelini hujumlarni aniqlash uchun tatbiq qilishgan. Ular *DNN* (*Deep Neural Network*) asosidagi modelni ishlab chiqib, KDD va CICIDS2017 ma‘lumotlar to‘plamlarida sinovdan o‘tkazishgan. Natijalar shuni ko‘rsatdiki, chuqur neyron tarmoqlar an‘anaviy ML modellariga nisbatan ancha yuqori aniqlikka ega [3].

Shunga o‘xshash tarzda, Yin(2017) o‘z tadqiqotida *RNN* va *LSTM* arxitekturalaridan foydalanib, tarmoq trafikidagi vaqt bo‘yicha ketma-ketliklarni tahlil qilishgan. Ularning modeli hujum sodir bo‘lishidan oldin paydo bo‘ladigan xatti-harakat naqshlarini aniqlay olgan. Natijalar *RNN* va *LSTM* yondashuvlari IDS tizimlari uchun istiqbolli yo‘nalish ekanini ko‘rsatgan [4].

Ferrag, Maglaras, Moschogiannis va Janicke (2020) esa chuqur o‘rganish asosidagi tarmoq hujumlarini aniqlash modellarini keng qamrovli tarzda tahlil qilishgan. Ular turli datasetlar, arxitekturalar (*CNN*, *RNN*, *Autoencoder*, *GAN*) va baholash metrikalarini solishtirib, sohadagi dolzarb muammolar — ma‘lumotlar yetishmasligi, balanssiz sinflar va real-vaqt aniqlashning murakkabligi haqida xulosalar bergan [5].

Khraisat(2019) tomonidan yozilgan “Survey of Intrusion Detection Systems” maqolasi esa bu yo‘nalishdagi eng mashhur tahliliy ishlar qatoriga kiradi. Ular hujumlarni aniqlash tizimlarini ikki asosiy turga — *signaturali* va *anomaliya asosidagi* tizimlarga ajratgan. Shuningdek, ishlatiladigan datasetlar va modellarning kuchli va zaif tomonlarini umumlashtirgan. Ularning tadqiqotida ayniqsa anomaliya asosidagi usullar yangi hujumlarni aniqlashda samaraliroq ekani ta‘kidlangan [6].

Yana bir muhim ish — Sarhan, Layeghy va Portmann (2021) tomonidan amalga oshirilgan tadqiqotdir. Ular ML asosidagi IDS modellarining umumlashuv (generalization) va tushuntirilish (explainability) xususiyatlarini o'rganishgan. Model turli datasetlarda sinovdan o'tkazilganda qanday o'zgarishlar bo'lishini tahlil qilishgan. Bu ish real tarmoq muhitida sun'iy intellekt modelining ishlash barqarorligini o'rganishda muhim bosqich hisoblanadi [7].

Umuman olganda, mavjud adabiyotlarni tahlil qilish quyidagi xulosalarni beradi:

- Sun'iy intellekt algoritmlari, ayniqsa chuqur o'rganish modellarining samaradorligi an'anaviy imzoli IDS tizimlariga nisbatan ancha yuqori.

- Anomaliya asosidagi yondashuvlar yangi turdagi hujumlarni aniqlashda samaraliroq, ammo *false positive* darajasi balandroq.

- Amaliy jihatdan qaraganda, modellarni real tarmoq sharoitida joylashtirishda (deploy qilishda) resurs sarfi va ishlash tezligi muhim omil bo'lib qolmoqda.

- Kelgusida GNN (Graph Neural Network) va Transformer kabi modellar tarmoq strukturasini yanada chuqurroq tahlil qilish imkonini beradi.

- Modelning *izohlanishi (explainability)* va *adversarial barqarorligi* hozirgi zamonaviy tadqiqotlarda asosiy yo'nalishlardan biri sifatida qaralmoqda.

Shu tarzda, adabiyotlar tahlili shuni ko'rsatadiki, tarmoq hujumlarini aniqlashda sun'iy intellekt algoritmlari katta imkoniyatlarga ega, ammo ularni real tizimlarda tatbiq qilishda hali ham ko'plab muammolar mavjud.

IV. NATIJA

Tarmoq hujumlarini aniqlashda sun'iy intellekt algoritmlarini qo'llash zamonaviy kiberxavfsizlik tizimlarining eng muhim yo'nalishlaridan biri bo'lib, bu sohada klassik imzo asosidagi yondashuvlardan boshlab, anomaliya asosidagi ilg'or o'rganish modellarigacha bo'lgan keng spektrdagi metodlar qo'llaniladi. Amalda tarmoq hujumlarini aniqlash tizimlari (IDS – Intrusion Detection Systems) ikki asosiy turga ajratiladi: *signaturali tizimlar* va *anomaliya asosidagi tizimlar*. Har ikkisi bir-birini to'ldiruvchi yondashuvlar hisoblanadi, biroq ularning ishlash mexanizmi, aniqlash usuli va qo'llaniladigan algoritmlari keskin farq qiladi.

Signaturali tizimlar, eng avvalo, *ma'lum hujum turlarining imzolari (patterns)* yoki *qoida to'plamlari* asosida ishlaydi. Ular avvaldan ma'lum bo'lgan zararli faoliyatni aniqlash uchun shablonlarni solishtiradi. Bunday tizimlarda hujumning aniqlanishi “moslik” tamoyiliga asoslanadi, ya'ni tarmoqda kuzatilayotgan trafikdagi baytlar yoki paketlarning xususiyatlari ma'lum hujum imzosiga mos kelsa, tizim ogohlantirish beradi. Bu turdagi yondashuv Snort, Suricata, Bro/Zeek kabi an'anaviy IDSlarda keng tarqalgan. Ammo signaturali tizimlarning asosiy kamchiligi — ular noma'lum yoki yangi hujumlarni aniqlay olmaydi, chunki har bir yangi hujum uchun alohida imzo qo'lda yozilishi kerak.

Sun'iy intellektning kirib kelishi bilan signaturali tizimlar mashinani o'rganish (Machine Learning) modellari yordamida avtomatik ravishda qoidalar hosil qilish va mavjud imzolarni optimallashtirish imkoniga ega bo'ldi. Masalan, *Decision Tree*, *Random Forest* yoki *Support Vector Machine (SVM)* kabi algoritmlar avvalgi trafikdan

o'rganilgan naqshlar asosida yangi kelayotgan ma'lumotlarni avtomatik tasniflay oladi. Shunday qilib, model oldindan belgilangan qoida emas, balki o'zi o'rganib chiqqan naqshlar asosida qaror chiqaradi. Misol uchun, *Random Forest* algoritmi turli atributlarga asoslanib, hujum ehtimolini baholaydi va aniq hujum turi bilan bog'laydi. *Naive Bayes* esa ehtimollik yondashuvi orqali ma'lum bir trafik segmentining hujum bo'lish ehtimolini hisoblaydi. Shu tarzda sun'iy intellekt signaturali tizimlarning imkoniyatlarini kengaytirib, qoida yaratish jarayonini avtomatlashtirgan.

Biroq bu yondashuv hali ham "oldindan belgilangan bilimlar"ga tayanadi. Shuning uchun zamonaviy kiberxavfsizlikda asosiy e'tibor *anomaliya asosidagi (Anomaly-based Detection)* tizimlarga qaratilmoqda. Bunday tizimlarda maqsad — avval "normal" tarmoq faoliyatini o'rganib olish va undan chetga chiqqan xatti-harakatlarni anomaliya sifatida aniqlashdir. Bu yondashuvning eng katta ustunligi — u noma'lum, ilgari uchramagan hujumlarni (zero-day attacks) aniqlay oladi.

Anomaliya aniqlashda ishlatiladigan sun'iy intellekt algoritmlari turlicha bo'lishi mumkin. Masalan, unsupervayzlangan (belgilarsiz) usullarda *k-means* yoki *DBSCAN* klasterlash algoritmlari ishlatiladi. Ular ma'lumotlar to'plamini guruhlariga ajratadi, va har qanday "klaster tashqarisidagi" element — potentsial hujum deb hisoblanadi.

So'nggi yillarda *chuqur o'rganish (Deep Learning)* asosidagi anomaliya aniqlash modellarining samaradorligi yuqori bo'lib chiqdi. Autoencoder modellar, masalan, normal trafikni o'rganib, keyinchalik uni qayta tiklaydi. Agar model yangi trafikni qayta tiklashda katta xatolik qilsa — bu hujum belgisi bo'lishi mumkin. Shuningdek, *RNN (Recurrent Neural Network)*, *LSTM (Long Short-Term Memory)*, *GRU (Gated Recurrent Unit)* tarmoqlari vaqt bo'yicha trafik oqimini tahlil qiladi va vaqt ketma-ketligida odatdagidan chetga chiqishlarni aniqlaydi. Bunday modellar DoS, DDoS yoki port scanning kabi hujumlarda vaqt ketma-ketlikdagi naqshlarni samarali aniqlaydi.

Yana bir istiqbolli yo'nalish — *Graf Neyron Tarmoqlari (Graph Neural Networks — GNN)* asosidagi yondashuvdir. Tarmoqdagi tugunlar va ularning o'zaro aloqalari graf shaklida tasvirlanadi, va model aynan shu graf tuzilmasini o'rganadi. Masalan, har bir tugun — bu IP manzil, har bir qirra (edge) esa ma'lumot oqimi. Agar ma'lum bir tugun yoki bog'lanish odatdagidan farqli harakat qilsa, model bu aloqani hujum sifatida baholaydi. Bu usul tarqatilgan (distributed) yoki ko'p nuqtali (multi-source) hujumlarni aniqlashda juda foydali.

Ba'zi tadqiqotlarda Generativ Adversarial Network (GAN) modellaridan ham foydalanilgan. GAN yordamida anomaliya namunalarini yaratish (data augmentation) yoki mavjud modelni hujumlarga nisbatan barqarorroq qilish (adversarial training) mumkin. Shu bilan birga, ensemble learning yondashuvi — ya'ni bir nechta modelni birlashtirish (masalan, *Random Forest + Autoencoder* yoki *CNN + LSTM* kombinatsiyasi) — yuqori aniqlik va barqarorlikni ta'minlaydi.

Sun'iy intellekt algoritmlarini amaliyotda qo'llashda bir nechta muhim bosqich mavjud: birinchidan, *ma'lumotlarni to'plash va oldindan ishlov berish*, ikkinchidan, *modelni o'qitish*, uchinchidan esa *baholash*. Tadqiqotchilar, masalan, CICIDS2017 va UNSW-NB15 datasetlarida aynan shu metrikalar asosida AI modellarining samaradorligini sinovdan o'tkazishgan.

AI asosidagi tizimlarning amaliy qiymati ularning real-vaqt aniqlash imkoniyatiga ham bog‘liq. Masalan, CNN-LSTM modellarini TensorFlow yoki PyTorch asosida ishlab chiqilgan IDS tizimlariga joylashtirish mumkin, bunda model o‘qitilgandan so‘ng real trafik oqimida inference (bashorat) bajaradi. Biroq, real vaqtda ishlov berish jarayonida kechikish (latency) muammosi paydo bo‘lishi mumkin, shuning uchun modellarning soddalashtirilgan (lightweight) versiyalarini ishlab chiqish dolzarb.

AI modellarining samaradorligini oshirish uchun *explainable AI (XAI)* yondashuvlari ham kiritilmoqda. Masalan, SHAP yoki LIME texnikalari modelning qaror chiqarish sabablarini tushuntirishga yordam beradi. Bu kiberxavfsizlik operatorlari uchun juda muhim, chunki ular modelning “nima sababdan bu trafikni hujum deb baholaganini” bilishga intilishadi.

Shu bilan birga, har qanday algoritmnining kuchli va zaif tomonlari mavjud. Quyidagi jadvalda signaturali va anomaliya asosidagi yondashuvlarning asosiy afzalliklari va kamchiliklari keltirilgan(1-jadval):

1-jadval

Yondashuv turi	Afzalliklari	Kamchiliklari
Signaturali (imzo asosida)	Yuqori aniqlik; aniq hujum turlarini tez aniqlaydi; noto‘g‘ri ogohlantirishlar kam	Yangi (zero-day) hujumlarni aniqlay olmaydi; imzolarni doim yangilab turish kerak
Supervayzlangan ML (SVM, RF, XGBoost)	Belgilangan ma’lumotda yuqori natija; tez o‘rganadi; interpretatsiya oson	Belgilangan ma’lumot zarur; imbalanced sinflarda samarasi past
Anomaliya asosidagi (Autoencoder, Isolation Forest)	Yangi hujumlarni aniqlay oladi; ma’lumot belgilarini talab qilmaydi	Ko‘p false positive beradi; threshold tanlash muhim
Chuqur o‘rganish (CNN, LSTM, Transformer)	Murakkab naqshlarni o‘rganadi; yuqori aniqlik; zero-day hujumlarga sezgir	Hisoblash resursi katta; tushuntirish (explainability) murakkab
GNN asosidagi yondashuvlar	Tarqatilgan hujumlarni aniqlaydi; tarmoq tuzilmasini hisobga oladi	Grafni qurish murakkab; katta xotira talab qiladi
Ensemble modellar	Barqarorlik va yuqori aniqlik; turli model xatolarini muvozanatlashtiradi	Implementatsiya murakkab; ishlash tezligi past bo‘lishi mumkin

V. XULOSA VA TAVSIYALAR

Xulosa qilib aytganda, tarmoq hujumlarini aniqlashda sun‘iy intellekt algoritmlarini qo‘llash an‘anaviy signaturali yondashuvlarga nisbatan sezilarli ustunlikka ega. Mashinani o‘rganish va chuqur o‘rganish modellarining tahlili shuni ko‘rsatdiki, ular katta hajmdagi tarmoq trafikini avtomatik tahlil qilish, murakkab naqshlarni aniqlash va yangi turdagi hujumlarni (zero-day attacks) aniqlash imkonini beradi.

Signaturali usullar (masalan, Decision Tree, Random Forest, SVM) aniq belgilangan hujumlarni aniqlashda samarali bo‘lsa, anomaliya asosidagi yondashuvlar

(Autoencoder, LSTM, Isolation Forest va boshqalar) noma'lum hujumlarni topishda afzal natijalar beradi. Shunga qaramay, anomaliya modellarida noto'g'ri ogohlantirishlar (false positive) darajasi yuqoriligi amaliy qo'llashda muammo tug'diradi. Shu bois, eng yaxshi natija **gibrid tizimlar** — ya'ni signaturali va anomaliya yondashuvlarini birlashtirgan modellar orqali olinadi.

Sun'iy intellektga asoslangan tizimlarning samaradorligi ma'lumotlar to'plamining sifatiga, xususiyatlarni to'g'ri tanlashga va modelning umumlashuv (generalization) qobiliyatiga bevosita bog'liq. Ayniqsa, chuqur o'rganish modellarining vaqt ketma-ketlikdagi hujumlarni tahlil qilishdagi aniqligi yuqori natijalarni ko'rsatdi.

Tavsiya sifatida, kelgusida tarmoq hujumlarini aniqlashda **graf-neyral tarmoqlar (GNN), Transformer** va **explainable AI** texnologiyalaridan foydalanish muhim yo'nalish bo'lib qoladi. Shuningdek, real-vaqt (real-time) ishlashni ta'minlash uchun yengil (lightweight) modellarni ishlab chiqish, modelni muntazam yangilab turish va ma'lumotlar xavfsizligini oshirish zarur.

FOYDALANILGAN ADABIYOTLAR

- [1].Khraisat, A., Gondal, I., Vamplew, P., & Kamruzzaman, J. (2019). *Survey of intrusion detection systems: techniques, datasets and challenges*. Security and Communication Networks, 2019.
- [2].Buczak, A. L., & Guven, E. (2016). *A survey of data mining and machine learning methods for cyber security intrusion detection*. IEEE Communications Surveys & Tutorials, 18(2), 1153–1176.
- [3].Vinayakumar, R., Alazab, M., Soman, K. P., Poornachandran, P., Al-Nemrat, A., & Venkatraman, S. (2019). *Deep Learning Approach for Intelligent Intrusion Detection System*. IEEE Access, 7, 41525–41550.
- [4].Yin, C., Zhu, Y., Fei, J., & He, X. (2017). *A Deep Learning Approach for Intrusion Detection Using Recurrent Neural Networks*. IEEE Access, 5, 21954–21961.
- [5].Ferrag, M. A., Maglaras, L., Moschogiannis, S., & Janicke, H. (2020). *Deep learning for cyber security intrusion detection: Approaches, datasets, and comparative study*. Journal of Information Security and Applications, 50, 102419.
- [6].Liao, H.-J., Lin, C.-H. R., Lin, Y.-C., & Tung, K.-Y. (2013). *Intrusion detection system: A comprehensive review*. Journal of Network and Computer Applications, 36(1), 16–24.
- [7].Sarhan, M., Layeghy, S., & Portmann, M. (2021). *Evaluating standard feature sets towards increased generalisability and explainability of ML-based network intrusion detection*. Journal of Information Security and Applications, 63, 103045.