

DEEP LEARNING У СИСТЕМАХ ВИЯВЛЕННЯ ВТОРГНЕНЬ: АНАЛІЗ АРХІТЕКТУР, ПЕРЕВАГ ТА ОБМЕЖЕНЬ

Рихва Володимир,

Аспірант

volodymyr.rykhva@hneu.net

Солодовник Ганна

к.т.н., доцент

ganna.solodovnyk@hneu.net

кафедри кібербезпеки та інформаційних технологій

Харківський національний економічний

університет імені Семена Кузнеця, Україна

У сучасному цифровому середовищі зростає кількість та складність кібератак, що ставить під загрозу безпеку комп'ютерних мереж. Традиційні системи виявлення вторгнень (IDS), які базуються на сигнатурному або евристичному аналізі, часто виявляються неефективними проти нових або модифікованих атак. У зв'язку з цим методи глибокого навчання (Deep Learning) набувають все більшого значення у розробці IDS, оскільки вони здатні автоматично виявляти складні шаблони та аномалії у мережевому трафіку [1].

Архітектури глибокого навчання у IDS

Deep Learning - це підмножина машинного навчання, яка в першу чергу фокусується на імітації здатності людського мозку навчатися та обробляти інформацію [2]. У світі штучного інтелекту (ШІ), що стрімко розвивається, глибоке навчання стало революційною технологією, яка впливає практично на всі сфери - від охорони здоров'я до автономних систем. Щоб досягти такої здатності до навчання та обробки інформації, глибоке навчання спирається на складну мережу взаємопов'язаних нейронів, які називаються штучними нейронними мережами (ШНМ). Використовуючи потужність ШНМ та їхню здатність автоматично адаптуватися і вдосконалюватися з часом, алгоритми глибокого навчання можуть виявляти складні закономірності, витягувати значущі ідеї та робити прогнози з надзвичайною точністю.

ШНМ складається з вхідного шару, вихідного шару та декількох прихованих шарів між ними. Вхідний шар приймає дані, а вихідний надає

результат роботи моделі. Якщо кількість прихованих шарів дорівнює одному, то нейронна мережа називається неглибокою, якщо більше – глибокою (Рис.1). Шари містять в собі вузли, які є штучними нейронами.

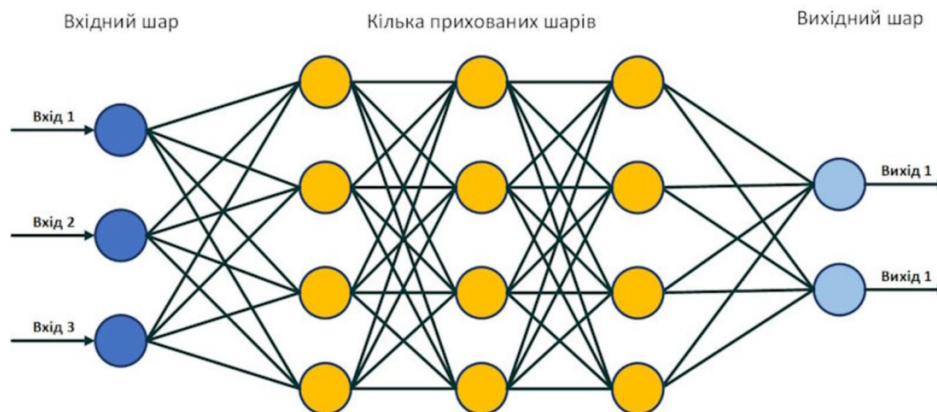


Рис.1 Архітектура глибокої нейронної мережі

На вхідний шар надходить інформація. Вузли (нейрони) цього шару обробляють вхідні сигнали, аналізують їх або класифікують. Потім дані передаються на наступний рівень. Кількість входів визначає кількість функцій, які використовує нейронна мережа для прогнозування. Кількість виходів в нейронній мережі визначається кількістю прогнозів. Кількість прихованих шарів залежить від проблеми та архітектури штучної нейронної мережі [3].

Згорткові нейромережі (CNN) ефективно обробляють дані з локальними залежностями, що дозволяє виявляти просторові патерни у мережевому трафіку. Їх застосування в IDS дозволяє автоматично виділяти релевантні ознаки без необхідності ручного інжинірингу. Проте, CNN мають обмежену здатність до обробки часових залежностей, що може знижувати ефективність у виявленні атак, які розгортаються у часі [4].

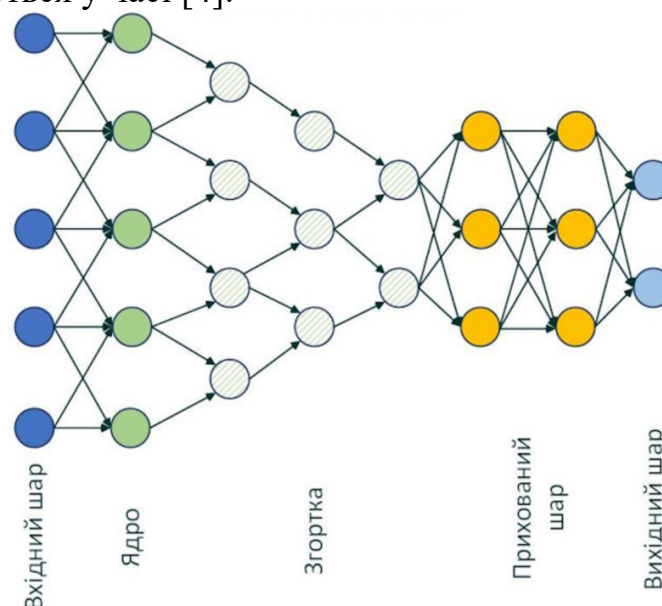


Рис.2 Архітектура CNN

Рекурентні нейронні мережі (RNN) та їх модифікація Long Short-Term Memory (LSTM) є ефективними інструментами для обробки послідовних даних, таких як мережевий трафік, де важливими є часові залежності (Рис. 3). Однак, традиційні RNN стикаються з проблемою згасання градієнта, що ускладнює навчання моделей на довгих послідовностях. LSTM-мережі були розроблені для подолання цієї проблеми шляхом введення спеціальних механізмів пам'яті, які дозволяють зберігати інформацію про попередні стани протягом тривалого часу. Це робить LSTM придатними для задач, що вимагають врахування довгострокових залежностей у даних. Проте, LSTM-мережі мають свої обмеження. Зокрема, вони можуть бути чутливими до дуже довгих послідовностей та вимагати значних обчислювальних ресурсів для ефективного навчання. Це пов'язано з їх складною архітектурою та необхідністю обробки великої кількості параметрів, що може призводити до високих витрат часу та ресурсів при навчанні моделей на великих наборах даних.

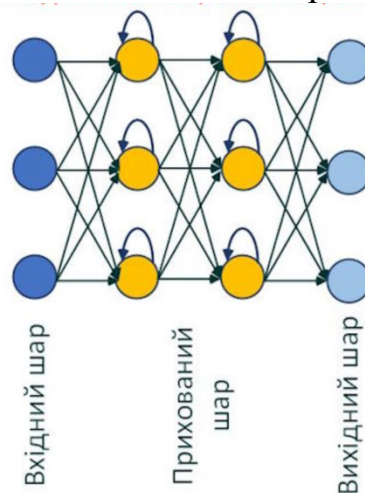


Рис.3 Архітектура RNN

Автоенкодери (АЕ) використовуються для виявлення аномалій шляхом навчання реконструкції нормального трафіку. Відхилення у реконструкції може свідчити про наявність аномалії або атаки. АЕ ефективні у виявленні невідомих атак, проте можуть мати високу кількість хибнопозитивних спрацювань [5].

Граф нейронної мережі (GNN) дозволяють моделювати мережевий трафік у вигляді графів, що забезпечує врахування топології мережі та взаємозв'язків між вузлами. Це особливо корисно для виявлення атак, які поширюються через мережу, таких як боти або розподілені атаки. [5].

Комбінація різних архітектур, наприклад CNN та LSTM, дозволяє об'єднати переваги обох підходів. CNN ефективно виділяють просторові ознаки, тоді як LSTM обробляють часові залежності. Такий підхід демонструє високу ефективність у виявленні складних атак у реальному часі [7].

Переваги та недоліки використання Deep Learning у IDS наведені у табл.1.

Таблиця 1 – Переваги та недоліки використання Deep Learning

Переваги	Недоліки
Автоматичне виділення ознак	Високі обчислювальні витрати
Висока точність виявлення як відомих, так і нових атак.	моделі часто розглядаються як “чорні ящики”, що ускладнює розуміння прийнятих ними рішень.
Адаптивність до нових загроз	Ефективне навчання моделей потребує великих обсягів якісних даних
Здатність до обробки складних шаблонів	Низька якість або незбалансованість даних може негативно впливати на продуктивність моделей.
Масштабованість (можуть бути ефективно розгорнуті на хмарних платформах та пристроях з обмеженими ресурсами)	Ризик перенавчання (overfitting)

Список використаних джерел

1. Столяр А. Л. Аналіз сучасних методів виявлення аномалій в комп'ютерних мережах. 2023, URL: <https://doi.org/10.18372/2073-4751.74.17888>.
2. Schmidhuber, J. (2015). Deep Learning in Neural Networks: An Overview. Neural Networks, 61, 85–117. <https://doi.org/10.1016/j.neunet.2014.09.003>
3. Sarker, I. H. (2021). Deep Learning: A Comprehensive Overview on Techniques, Taxonomy, Applications and Research Directions. SN Computer Science, 2(6), 420. <https://doi.org/10.1007/s42979-021-00815-1>
4. Kimanzi, R., Kimanga, P., Cherori, D., & Gikunda, P. K. Deep Learning Algorithms Used in Intrusion Detection Systems – A Review. arXiv preprint arXiv:2402.17020, 2024.
5. Singh, A., & Jang-Jaccard, J. Autoencoder-based Unsupervised Intrusion Detection using Multi-Scale Convolutional Recurrent Networks. arXiv preprint arXiv:2204.03779, 2022.
6. Zhong, M., Lin, M., Zhang, C., & Xu, Z. A survey on graph neural networks for intrusion detection systems: methods, trends and challenges. Computers & Security, 141, 103821, 2024.
7. Sinha, P., Sahu, D., Prakash, S., Yang, T., Rathore, R. S., & Pandey, V. K. A high performance hybrid LSTM CNN secure architecture for IoT environments using deep learning. Scientific Reports, 15, 9684, 2025.