

Bài báo nghiên cứu

**MÔ HÌNH HỌC SÂU LONG SHORT-TERM MEMORY
PHÁT HIỆN TẤN CÔNG DDOS****Phạm Trọng Huỳnh**

Trường Đại học Tài Nguyên và Môi trường Thành phố Hồ Chí Minh, Việt Nam

Tác giả liên hệ: Phạm Trọng Huỳnh – Email: pthuynh@hcmunre.edu.vn

Ngày nhận bài: 13-6-2024; ngày nhận bài sửa: 11-12-2024; ngày duyệt đăng: 23-01-2025

TÓM TẮT

Gần đây, các mối đe dọa tấn công Từ chối dịch vụ phân tán-Distributed Denial of Service (DDoS) đang trở nên phức tạp, tinh vi, gây ra thách thức cho các hệ thống bảo vệ thông thường. Việc phát hiện sớm các dấu hiệu tấn công rất quan trọng, để bảo vệ và chống lại các mối đe dọa tấn công. Nghiên cứu đề xuất sử dụng mô hình dựa trên kỹ thuật Học sâu Mạng bộ nhớ Dài-Ngắn - Long Short-Term Memory (LSTM). Kỹ thuật LSTM này gồm một số thuật toán lựa chọn và trích xuất đặc trưng, được tự động cập nhật trong quá trình huấn luyện. Với số lượng dữ liệu nhỏ, LSTM vẫn hoạt động nhanh và chính xác. Nghiên cứu đã tiến hành thử nghiệm trên tập dữ liệu CICDDoS2019 và kết quả cho thấy mô hình đạt được các chỉ số hiệu suất như sau: Độ chính xác (Accuracy) đạt 93%, độ chuẩn xác (Precision) đạt 96%, độ phủ (Recall) đạt 93% và điểm F1 (F1 Score) là 94%. Mục tiêu của nghiên cứu, đưa ra được một mô hình có khả năng xử lý dữ liệu chuỗi và lưu trữ thông tin học được lâu dài. Có thể tích hợp mô hình vào các hệ thống giám sát và bảo mật mạng, cải thiện khả năng phát hiện phản ứng với các mối đe dọa tấn công mạng ngày càng phức tạp.

Từ khóa: DDoS; học sâu; DoS; LSTM; học máy**1. Giới thiệu**

Tốc độ phát triển nhanh chóng của các dịch vụ thông qua mạng Internet như: giao dịch tài chính và ngân hàng, truyền thông, thương mại điện tử, mua sắm, thanh toán trực tuyến, chăm sóc sức khỏe và giáo dục.

Việc bảo vệ an toàn cho người dùng đang là một thách thức. Hiện có rất nhiều phương thức tấn công mạng nhằm mục đích phá hoại, trong đó hai phương pháp tấn công cơ bản nhất vẫn là tấn công từ chối dịch vụ (Denial of Service - DoS), tấn công nhằm ngăn chặn người dùng hợp pháp truy nhập các tài nguyên mạng và tấn công từ chối dịch vụ phân tán (Distributed Denial of Service – DDoS) là một dạng phát triển hơn của tấn công DoS (Kumar, 2020). Các phương pháp truyền thống để xác định rủi ro của các cuộc tấn công DDoS thường có độ chính xác thấp và phản ứng chậm (Dave et al., 2022). Để giải quyết vấn

Cite this article as: Phạm Trọng Huỳnh (2025). Long short-term memory deep learning model detecting DDoS attacks. *Ho Chi Minh City University of Education Journal of Science*, 22(2), 224-234.

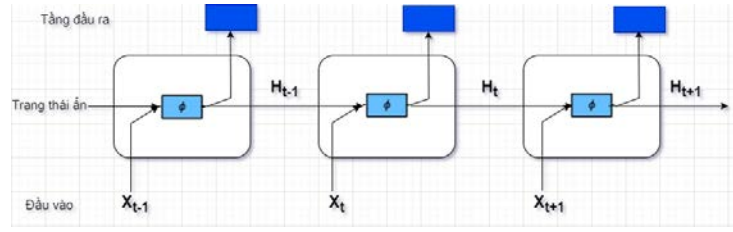
đề này, các chuyên gia an ninh mạng và học giả đã và đang nghiên cứu sử dụng Máy học - Machine Learning và Học sâu trong việc phát hiện DDoS (Dincalp, 2018). Các phương pháp ML và DL có tiềm năng vô cùng to lớn trong việc phát hiện các cuộc tấn công mạng bởi chúng có khả năng phân loại một cách chính xác và hiệu quả hơn (Fadlil, 2017). Hiện nay các phương pháp như Random Forest (RF), K-Nearest Neighbors (KNN) và Naive Bayes đang được sử dụng (Zheng et al., 2017). Đối với Học sâu, thường sử dụng các phương pháp như: Mạng nơ-ron nhân tạo -Artificial Neural Networks (ANN), Mạng nơ-ron sâu -Deep Neural Network (DNN) và Mạng nơ-ron hồi quy-Recurrent Neural Network (RNN) (Ahanger, 2017). Sau khi khảo sát và đánh giá nhiều phương pháp, nghiên cứu quyết định áp dụng kiến trúc Học sâu Mạng bộ nhớ Dài-Ngắn - Long short-Term Memory (LSTM), để thực hiện (Zahid et al., 2018). Kiến trúc Học sâu LSTM có thể thu thập được thông tin, dựa trên tính chất tương quan trong miền thời gian, đây là một đặc điểm quan trọng của lưu lượng DDoS, vì thế LSTM có thể giúp dự đoán hiệu quả lưu lượng mạng. Mục tiêu của bài báo, là xây dựng một mô hình kiến trúc Học sâu LSTM đạt được độ chính xác cao hơn so với kỹ thuật Học máy đã có, trong việc phân loại và dự báo trước các cuộc tấn công DDoS.

Trong thời gian gần đây, các công trình nghiên cứu trên tập dữ liệu mẫu CICDDoS2019 về dự đoán các cuộc tấn công DDoS cho kết quả chính xác khá cao. Mỗi công trình nghiên cứu ứng dụng một kỹ thuật với những thế mạnh riêng. Trong nghiên cứu của (Ahuja et al., 2021), kết quả của dự đoán đạt 95% .

2. Đối tượng và phương pháp nghiên cứu

Mạng hồi quy (RNN) là một dạng mạng nơ-ron được thiết kế để xử lý dữ liệu theo chuỗi thời gian (Brunswick, 2021). Ý tưởng chính của RNN là sử dụng một bộ nhớ để giữ lại thông tin từ các bước tính toán trước đó, từ đó có thể đưa ra dự đoán chính xác cho các bước tiếp theo. Kiến trúc của RNN có thể được biểu diễn dưới dạng một chuỗi các đơn vị hồi quy (Zhu et al., 2018). Mỗi đơn vị này kết nối với đơn vị trước đó, tạo thành một chu trình có hướng. Ở mỗi điểm thời gian, đơn vị hồi quy nhận đầu vào hiện tại, kết hợp nó với trạng thái ẩn từ bước trước đó. Sau đó, đơn vị này tạo ra một đầu ra và cập nhật trạng thái ẩn cho bước thời gian tiếp theo. Quá trình này lặp lại cho mỗi đầu vào trong chuỗi, cho phép mạng hồi quy thu thập thông tin về các mối quan hệ và mẫu theo thời gian.

Quá trình tính toán của một RNN tại ba bước thời gian liên tiếp. Tại bước thời gian t , sau khi nối đầu vào X_t với trạng thái ẩn $H_{(t-1)}$ tại bước thời gian trước được kết hợp và đưa vào một tầng kết nối đầy đủ với hàm kích hoạt ϕ . Đầu ra của tầng này là trạng thái ẩn H_t , tại bước thời gian t cũng là đầu vào cho tầng đầu ra O_t . Các tham số mô hình W_{xh} , W_{hh} , cùng với H_t được sử dụng để tính toán trạng thái ẩn $H_{(t+1)}$ tại bước thời gian tiếp theo $t+1$. Quá trình này giúp RNN hiểu và duy trì thông tin quan trọng qua các bước thời gian, làm cho nó phù hợp cho các nhiệm vụ đòi hỏi xử lý chuỗi dữ liệu như dự đoán chuỗi thời gian. Mô hình kiến trúc RNN như trong Hình 1.



Hình 1. Kiến trúc RNN

Để quản lí dữ liệu và đưa dữ liệu vào mô hình, mô hình LSTM sẽ sử dụng ba cổng được đặt tên là cổng quên, cổng đầu vào và cổng đầu ra. Các cổng này là thành phần chính của mô hình LSTM và chịu trách nhiệm cho việc kiểm soát chung của mô hình. Cổng quên xác định phần nào của thông tin cũ nên loại bỏ dựa trên trạng thái ẩn trước đó và giá trị hiện tại của dữ liệu đầu vào. Cổng đầu vào quyết định loại dữ liệu nào được phép nhập vào mạng dựa trên thông tin liên quan nên được giới thiệu cho LSTM của mạng, được gọi là trạng thái ô nhớ (cell state). Cổng đầu ra xác định trạng thái ẩn mới dựa trên trạng thái ô nhớ đã được cập nhật và giá trị hiện tại của dữ liệu đầu vào.

Công thức tính toán cho các cổng:

$$F_{(t)} = \sigma(w_f[h_{(t-1)}, X_t] + b_f) \quad (1)$$

$$I_{(t)} = \sigma(w_i[[h_{(t-1)}, X_t] + b_i]) \quad (2)$$

$$O_{(t)} = \sigma(W_0 * [h_{(t-1)}, X_t] + b_0) \quad (3)$$

$$f(x) = \frac{1}{1+e^{-x}} \quad (4)$$

$$\tanh(x) = \frac{2}{1+e^{-2x}} - 1 \quad (5)$$

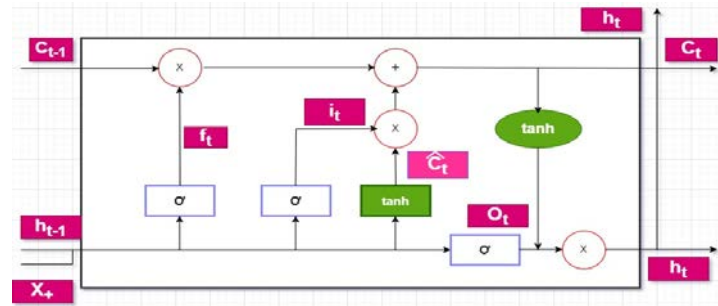
$$\hat{C}_t = \tanh(w_c, [h_{t-1}, X_t] + b_c) \quad (6)$$

$$c_t = F_t \cdot C_{(t-1)} + I_t \cdot \hat{C}_t \quad (7)$$

$$h_t = O_t \cdot \tanh(C_t) \quad (8)$$

trong đó:

$h_{(t-1)}$ là trạng thái ẩn trước đó, $x_{(t)}$ là giá trị hiện tại của dữ liệu đầu vào, $F_{(t)}$, $I_{(t)}$ và $O_{(t)}$ lần lượt là giá trị của các cổng quên, đầu vào và đầu ra, và W và b là các trọng số và bias tương ứng (1). Hàm sigmoid (σ) được sử dụng để lấy thông tin từ đầu vào gần nhất cũng như lớp ẩn trước đó. Phạm vi của hàm sigmoid (σ) dao động từ 0 đến 1 và phạm vi của hàm tanh dao động từ -1 đến 1. Nếu giá trị của hàm sigmoid gần bằng 1 thì nó giữ lại dữ liệu, còn nếu gần bằng 0 thì nó loại bỏ dữ liệu. Mô hình kiến trúc LSTM như trong Hình 2.



Hình 2. Kiến trúc LSTM

3. Kết quả và thảo luận

3.1. Môi trường thử nghiệm

Trong thử nghiệm này được chạy trên Python phiên bản 3.9 cùng với các thư viện học máy như Tensorflow và Sklearn và các thư viện hỗ trợ khác. Máy tính được sử dụng có bộ vi xử lý 8 nhân và 16 luồng, với bộ nhớ RAM 32GB. Ngoài ra, cũng sử dụng GPU để tăng tốc độ tính toán, với GPU NVIDIA GeForce RTX 3080.

3.2. Bộ dữ liệu

Trong nghiên cứu này sử dụng tập dữ liệu mẫu CICDDoS2019. Bộ dữ liệu “DDoS Evaluation Dataset (CICDDoS2019)” được Viện An ninh mạng Canada công bố vào ngày 31 tháng 10 năm 2019, bộ dữ liệu có 225.745 mẫu và 85 cột được lưu dưới định dạng CSV, bộ dữ liệu này cung cấp một tập hợp đa dạng các thông tin về các loại lưu lượng mạng và các biểu hiện của cuộc tấn công (cic/datasets/ddos-2019). Có nhiều nghiên cứu đã được triển khai trên tập dữ liệu này với những khía cạnh khác nhau, trong đó có một số mô hình học máy như Mạng Nơ-ron tích chập -CNNs (Doriguzzi-Corin., 2020), đạt được kết quả nhất định. Tuy vậy, dữ liệu của các cuộc tấn công DDoS dưới dạng chuỗi thời gian liên tục, việc áp dụng mô hình CNNs sẽ không hiệu quả. Vì mô hình này không tối ưu hóa cho dữ liệu tuần tự, thường gây mất thông tin về thứ tự thời gian do các phép gộp. Đặc biệt CNNs sẽ không nắm bắt tốt các mối tương quan dài hạn trong dữ liệu chuỗi thời gian. Ngoài ra khi sử dụng CNNs cho chuỗi thời gian, việc tiền xử lý dữ liệu để chuyển đổi nó thành định dạng phù hợp có thể phức tạp và tốn nhiều thời gian.

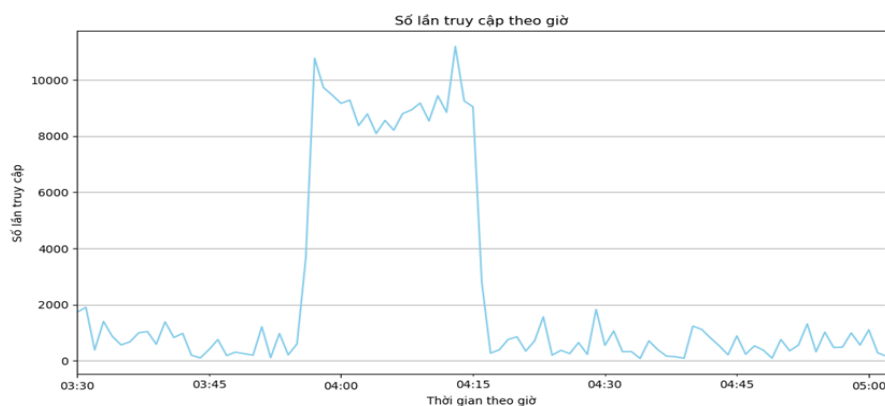
Quá trình phân loại dữ liệu đã xác định được 128.027 mẫu thuộc nhóm DDoS và 97.718 mẫu thuộc nhóm BENIGN. Sự phân phối của các nhãn trong tập dữ liệu tương đối cân bằng giữa hai nhóm, với khoảng 43,3% mẫu được xác định là không gây hại và 56,7% mẫu được xác định là tấn công DDoS.

Việc cân bằng giữa các nhãn trong tập dữ liệu là một yếu tố quan trọng khi huấn luyện mô hình. Điều này giúp đảm bảo rằng mô hình không chỉ hiểu được các cuộc tấn công mạng, mà còn có khả năng phân biệt chúng với các hoạt động hợp pháp một cách hiệu quả. Trong thực tế có nhiều yếu tố ảnh hưởng tới hiệu suất của mô hình dẫn tới quá trình dự đoán chưa được như kì vọng.

3.3. Trực quan hoá dữ liệu

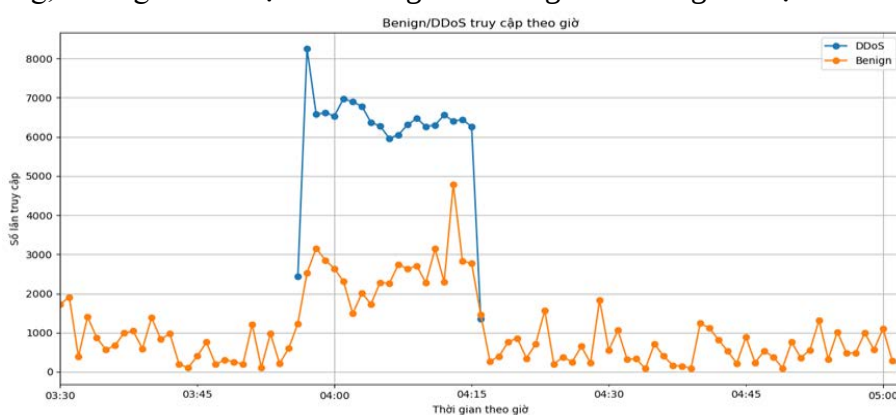
Việc trực quan hóa sự biến đổi của các đặc điểm khác nhau theo thời gian trên tập dữ liệu, để nhận biết các mẫu tấn công, thông qua lưu lượng truy cập bình thường và bất thường. Bằng cách này, có thể quan sát sự thay đổi của các biến và hiểu rõ hơn về các xu hướng, chu kỳ của dữ liệu theo thời gian thực. Sự biến đổi của các đặc điểm theo thời gian có thể cung cấp thông tin quan trọng trong quá trình phân tích dữ liệu và huấn luyện mô hình. Dữ liệu theo từng thời điểm (timestamp) trên tập dữ liệu DDoS SDN, được trực quan hóa, bằng cách nhóm dữ liệu theo thời điểm (timestamp) và tính số lượng dòng trong mỗi nhóm và hiển thị

dưới dạng biểu đồ để có thể quan sát sự phân phối của dữ liệu qua thời gian và nhận biết các biểu hiện đặc biệt trong dữ liệu.



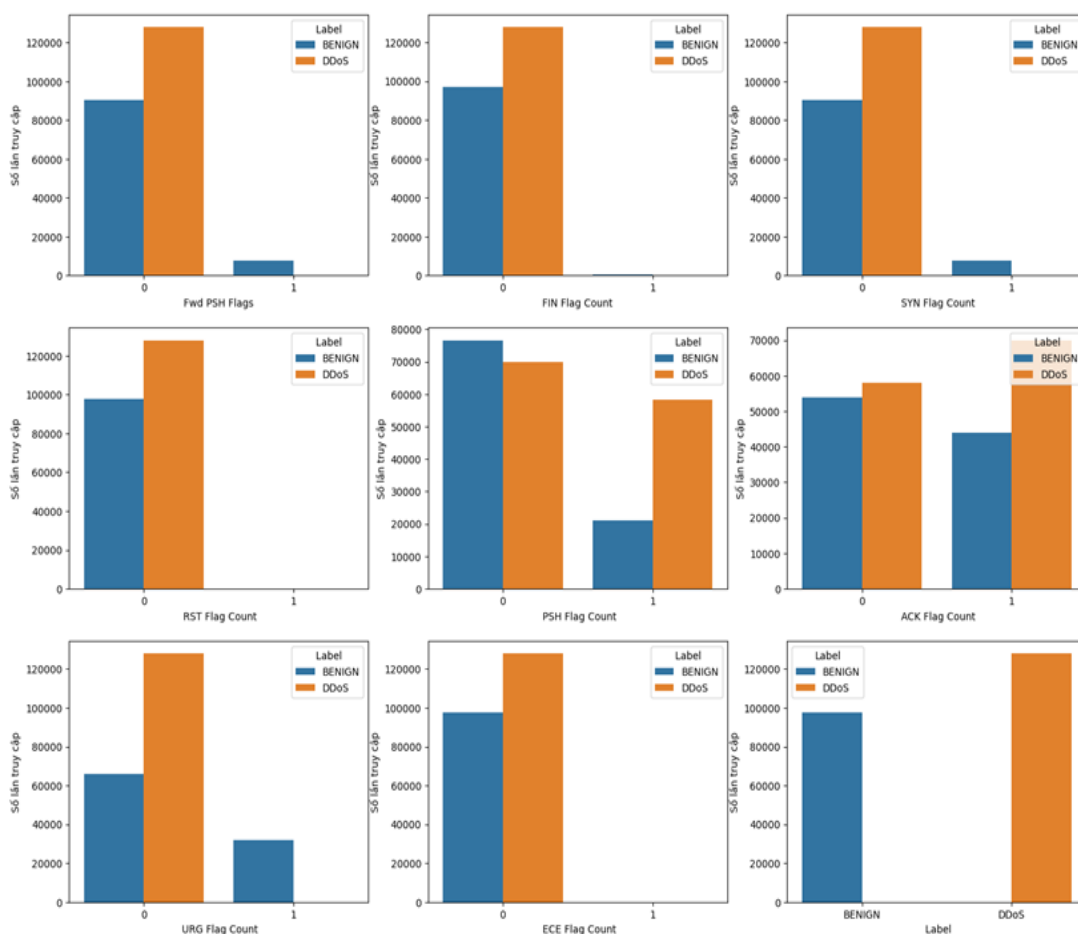
Hình 3. Biểu đồ biểu diễn dữ liệu theo từng thời điểm

Ngoài ra để so sánh số lượng kết nối của các cuộc tấn công DDoS và các kết nối không phải là cuộc tấn công DDoS (Benign) qua thời gian trên tập dữ liệu, bằng cách nhóm các dòng dữ liệu có nhãn “BENIGN” theo thời điểm và tính số lượng trong mỗi nhóm và hiển thị ra dưới dạng biểu đồ, qua đó ta có thể nhận biết sự biến động của các cuộc tấn công và các hoạt động không phải là tấn công trong hệ thống. Nhằm hỗ trợ việc phát hiện sớm các cuộc tấn công, đánh giá mức độ ảnh hưởng của chúng theo thời gian thực.



Hình 4. Biểu đồ thể hiện số lượng các kết nối tấn công và bình thường theo thời gian thực

Trong quá trình phân tích dữ liệu, đã tiến hành một loạt các công việc để khám phá sâu hơn về các đặc điểm và mối quan hệ giữa chúng trong tập dữ liệu DDoS SDN. Đầu tiên là đã xác định nhãn phân phối của các đặc điểm phân loại để hiểu rõ hơn về sự phân bố của chúng trong tập dữ liệu. Bước tiếp theo tiến hành phân tích và giải thích các thuộc tính TCP flags như PSH, FIN, SYN, RST, ACK, URG và ECE, nhằm hiểu được thông điệp và tác động của chúng trong mạng. Cuối cùng, trực quan hóa biến số nhóm theo nhãn trong dữ liệu, giúp nhận biết sự phân bố và mối quan hệ giữa các nhãn trong tập dữ liệu. Qua đó cung cấp thông tin của tập dữ liệu một cách chi tiết, rõ ràng, hỗ trợ quá trình phân tích và đưa ra các quyết định trong quá trình thực thi hệ thống cũng như bảo mật và quản lý mạng.



Hình 5. Biểu đồ phân nhóm các nhãn trong tập dữ liệu

3.4. Chuẩn hóa dữ liệu

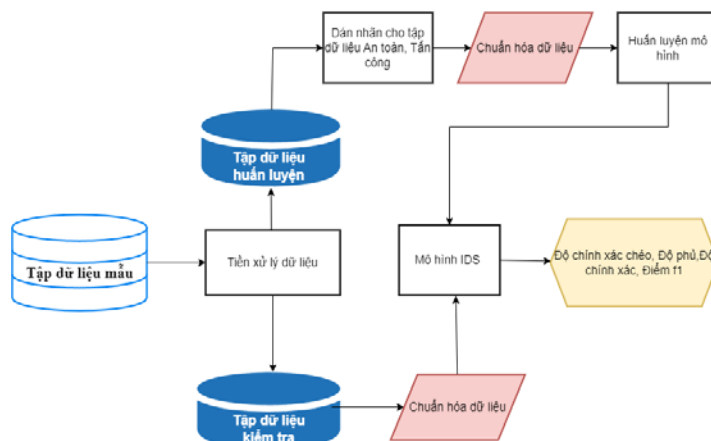
Quá trình làm sạch và biến đổi dữ liệu là một phần quan trọng của quy trình chuẩn bị dữ liệu. Đầu tiên, tiến hành loại bỏ các đặc trưng không đổi, những đặc trưng không mang lại thông tin hữu ích vì chúng không thay đổi qua các mẫu dữ liệu. Bước này giúp giảm kích thước của tập dữ liệu và loại bỏ nhiều không cần thiết.

Tiếp theo, tiến hành loại bỏ các đặc trưng dư thừa, tức là những đặc trưng có thể suy luận từ các đặc trưng khác. Điều này giúp giảm chi phí tính toán và tránh việc mô hình trở nên quá phức tạp. Trên tập dữ liệu này, đã loại bỏ các đặc trưng không đổi như Flow ID', 'Destination IP', 'Source IP', 'Destination Port', 'Source Port'. Do luồng Flow ID là duy nhất, bên cạnh đó khi thực hiện tấn công những IP này thường chạy ẩn danh hoặc là IP giả, vì vậy cần phải loại bỏ trường IP và Port. Dữ liệu quan sát cho thấy trường 'Protocol' gồm ba giao thức: UDP, TCP và POP3, trong đó TCP chiếm tới 85,5%. Hai giao thức còn lại xuất hiện với tần suất rất thấp nên gần như không đóng góp đáng kể vào việc dự đoán biến mục tiêu. Kết quả là tập dữ liệu hiện chỉ còn lại 225,745 hàng và 73 cột.

Bước tiếp theo, thực hiện mã hóa nhãn cho các đặc trưng phân loại bằng cách sử dụng LabelEncoder để chuyển đổi các nhãn sang dạng nhị phân. Sau đó, tiến hành phân tích mối tương quan giữa các đặc trưng trong tập dữ liệu để hiểu về mối quan hệ giữa chúng.

Cuối cùng, thực hiện quá trình lựa chọn đặc trưng bằng cách loại bỏ các cột có mối tương quan cao để giảm thiểu đa cộng tuyến và dư thừa trong tập dữ liệu. Việc này được lặp lại để xem xét mối tương quan giữa các đặc trưng và biến mục tiêu, nhằm giữ lại những đặc trưng có ý nghĩa và độc lập nhất. Quá đó giúp tinh chỉnh tập đặc trưng và cải thiện hiệu suất dự đoán của mô hình.

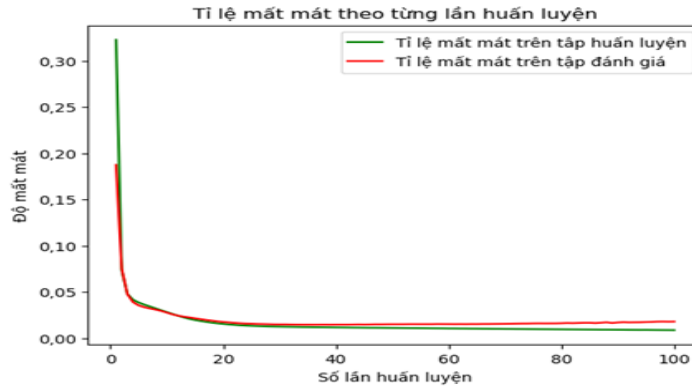
Quá trình huấn luyện mô hình LSTM trên dữ liệu đã được chuẩn hóa. Chia tập dữ liệu thành tập dữ liệu huấn luyện trainX và nhãn tương ứng trainY để điều chỉnh các trọng số của mô hình. Mô hình được huấn luyện qua 100 epochs. Để quản lý việc huấn luyện trên các lô dữ liệu nhỏ, kỹ thuật chia dữ liệu thành các batch có kích thước 1024 mẫu dữ liệu trong mỗi batch được thực hiện. Việc đánh giá hiệu suất của mô hình và tránh quá khớp, 20% của dữ liệu huấn luyện được sử dụng làm dữ liệu validation. Quá trình huấn luyện được hiển thị trên từng epoch, giúp theo dõi tiến triển của mô hình và hiểu rõ về sự biến thiên của các metric như loss và accuracy qua các vòng lặp. Cuối cùng, thông tin về quá trình huấn luyện được lưu trữ trong biến history_log, bao gồm các giá trị của các metric trong quá trình huấn luyện và validation sau mỗi epoch. Điều này giúp phân tích và đánh giá hiệu suất của mô hình sau khi huấn luyện hoàn tất. Quá trình thu thập, xử lý dữ liệu, huấn luyện mô hình được thể hiện thông qua mô hình tổng quát như trong Hình 7.



Hình 6. Mô hình LSTM đề xuất

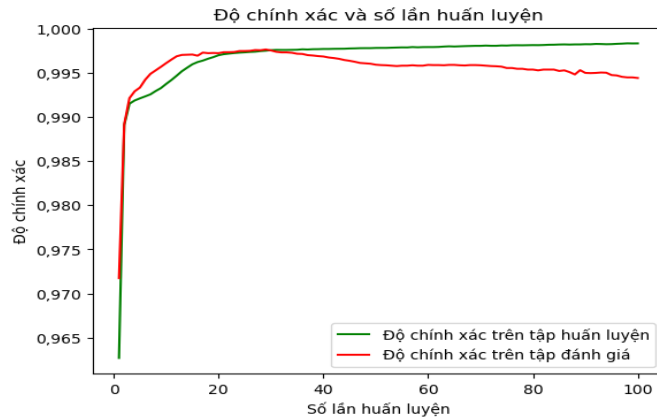
3.5. Đánh giá hiệu suất

Sự mất mát trong quá trình huấn luyện ban đầu khá cao và dần giảm qua các epoch, cho thấy mô hình đang học được. Sự mất mát trong quá trình validation cũng giảm đi, nhưng mức độ giảm này thấp hơn, cho thấy rằng mô hình khá tốt trong việc tổng quát hóa với dữ liệu mới chưa được sử dụng trong quá trình huấn luyện.



Hình 7. Biểu diễn giá trị Loss của mô hình

Bên cạnh đó độ chính xác trong quá trình học ban đầu ở mức trung bình, sau đó nhanh chóng tăng lên đạt được mức tốt nhất. Độ chính xác của quá trình validation cũng có xu hướng tương tự, mặc dù còn chậm hơn so với quá trình huấn luyện.



Hình 8. Biểu diễn giá trị Accuracy Score

Để đánh giá hiệu suất của các thuật toán học máy và học sâu, việc lựa chọn các tiêu chí đánh giá phù hợp là rất quan trọng. Nghiên cứu này, sử dụng các chỉ số đo độ chính xác (Accuracy - A), độ chính xác chéo (Precision - P), độ bao phủ (Re-call - R) và điểm F1 (F1-score) để đánh giá hiệu suất của mô hình. Công thức (9) tính toán độ chính xác. Giá trị nằm trong phạm vi từ 0 đến 1:

Precision đo lường tỉ lệ của các trường hợp dự đoán đúng trong số các trường hợp được dự đoán là tích cực. Nó được tính bằng cách chia số lượng các trường hợp dự đoán đúng (true positives) cho tổng số lượng các trường hợp được dự đoán là tích cực (true positives và false positives). Theo công thức:

$$\text{Precision} = \frac{\text{True Positive (TP)}}{\text{True Positive (TP)} + \text{False Positive (FP)}} \quad (8)$$

Recall: đo lường tỉ lệ của các trường hợp dự đoán đúng trong số các trường hợp thực sự tích cực. Nó được tính bằng cách chia số lượng các trường hợp dự đoán đúng (true positives) cho tổng số lượng các trường hợp thực sự tích cực (true positives và false negatives), được tính theo công thức:

$$\text{Recall} = \frac{\text{True Positive (TP)}}{\text{True Positive (TP)} + \text{False Negative (FN)}} \quad (9)$$

Accuracy: là đo lường tỉ lệ các dự đoán đúng (bao gồm cả dự đoán tích cực và dự đoán tiêu cực) so với tổng số lượng điểm dữ liệu. Được tính qua công thức:

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+PN} \quad (10)$$

F1 score là một giá trị trung bình điều hòa giữa Recall và Precision, được tính theo công thức sau:

$$\text{F1 Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (11)$$

Dựa vào các chỉ số đánh giá hiệu suất này, có thể đánh giá chất lượng của mô hình được đề xuất và đưa ra các quyết định phù hợp để cải thiện mô hình trong tương lai.

Bảng 1. Kết quả hiệu suất của mô hình trên tập dữ liệu Test

	Precision	Recall	F1-score
BENIGN	0,96	0,92	0,94
DDos	0,89	0,93	0,91
Accuracy			0,93
Macro avg	0,92	0,93	0,92
Weighted avg	0,93	0,93	0,93

3.6. Kết quả

Nghiên cứu này tập trung vào đánh giá hiệu suất của mô hình Long Short Term Memory (LSTM) thông qua việc sử dụng bộ dữ liệu DDoS SDN. Bộ dữ liệu này được chia thành hai phần: một phần được sử dụng để huấn luyện mô hình và một phần được dành cho kiểm thử. Các chỉ số đo hiệu suất như Accuracy, Recall, F1-score, Precision được sử dụng để thực hiện so sánh.

Trên tập dữ liệu CICDDoS2019 đã có nhiều công trình nghiên cứu với độ chính xác cao, tuy nhiên mỗi kỹ thuật có một đặc trưng riêng có của nó. Đối với kỹ thuật Optimized MLP-CNN Model, phù hợp để phát hiện các mẫu phức tạp từ dữ liệu không tuần tự. CNN giúp trích xuất đặc trưng không gian tốt, nhưng không mạnh trong xử lý chuỗi thời gian.

LSTM Model phù hợp hơn với các kịch bản yêu cầu phân tích chuỗi thời gian, đặc biệt là các cuộc tấn công DDoS có tính chất kéo dài hoặc lặp lại.

Kết quả của đánh giá hiệu suất mô hình trên tập kiểm thử được trình bày trong Bảng 2, cung cấp cái nhìn tổng quan về khả năng dự đoán của mô hình trong môi trường thực tế. Đóng vai trò quan trọng trong việc đánh giá.

Bảng 2. Kết quả hiệu suất của mô hình trên tập dữ liệu Test

Tham chiếu	Dataset	Accuracy	Precision	Recall	F1-score	Phương pháp
Long Short-Term Memory Model	CICDDoS2019	0,93	0,96	0,93	0,94	LSTM
(Ahuja et al., 2021).	CICDDoS2019	0,95	0,95	0,92	0,94	CNN
Thakkar et al., 2023)	CICDDoS2019	0,96	0,95	0,96	0,96	MLP-CN

4. Kết luận

Phát hiện kịp thời các cuộc tấn công DDoS từ xa là yếu tố quan trọng trong việc đảm bảo an ninh mạng và bảo vệ thông tin cá nhân. Nghiên cứu mới đề xuất sử dụng mô hình Học sâu Long Short Term Memory (LSTM) để dự báo các cuộc tấn công DDoS, và kết quả thử nghiệm cho thấy hiệu quả vượt trội so với các phương pháp học máy truyền thống. Mô hình LSTM đã chứng minh khả năng linh hoạt trong lựa chọn và trích xuất đặc trưng, và đạt được độ chính xác lên tới 93% trong việc phân loại các cuộc tấn công DDoS. Điều này vượt xa các mô hình truyền thống khác khi được đào tạo và đánh giá trên cùng tập dữ liệu. Sử dụng mô hình LSTM để phát hiện các cuộc tấn công DDoS có thể làm nền tảng cho nghiên cứu về phát hiện xâm nhập với độ chính xác cao. Kết hợp mô hình LSTM vào các mạng dựa trên phần mềm có thể là lựa chọn tốt, và tiếp tục nghiên cứu về tính năng học tăng cường bằng cách sử dụng đường dẫn lưu lượng mạng có thể cập nhật với các loại tấn công mới. Đề xuất này mang lại một giải pháp mới để giải quyết các vấn đề bảo mật mạng hiện nay, với khả năng phát hiện và ứng phó với các cuộc tấn công DDoS một cách hiệu quả, nhờ vào hiệu suất cao của mô hình LSTM đã được chứng minh.

❖ **Tuyên bố về quyền lợi:** Tác giả xác nhận hoàn toàn không có xung đột về quyền lợi.

TÀI LIỆU THAM KHẢO

- Ahanger, T. A. (2017). An effective approach of detecting DDoS using artificial neural networks. In *Proceedings of the 2017 International Conference on Wireless Communications, Signal Processing and Networking, Chennai, India*, (pp.22-24, March 2017). IEEE: Piscataway Township, NJ, USA.
- Ahuja, N., Singal, G., Mukhopadhyay, D., & Kumar, N. (2021). Automated DDoS attack detection in software defined networking. *Journal of Network and Computer Applications*, 187, Article 103108.
- Alzahrani, S., & Hong, L. (2018). Detection of distributed denial of service (DDoS) attacks using artificial intelligence on cloud. In *2018 IEEE World Congress on Services (SERVICES)* (pp.35-36). San Francisco, CA.
- Deepak Kumar, R. K. Pateriya, Rajeev Kumar Gupta, Vasudev Dehalwar, & Ashutosh Sharma (2019). *Computer Science and Engineering Department, Maula-na Azad National Institute of Technology, Bhopal, India*, 462003.
- Dincalp. (2018). Anomaly based distributed denial of service attack detection and prevention with machine learning. In *Proceedings of the 2nd International Symposium on Multidisciplinary Studies and Innovative Technologies* (pp.19-21, October 2018). Ankara, Turkey.

- Divyang Dave, Meet Kava, R. K. Gupta, & Kaushal Shah. (2022). Deep Learning approach for Intrusion Detection System. *IEEE International Conference on Technology, Research, and Innovation for Betterment of Society (TRIBES)*, 2022. <https://doi.org/10.1109/TRIBES52498.2021.9751643>
- Doriguzzi-Corin, R., Millar, S., Scott-Hayward, S., Martinez-del-Rincon, J., & Siracusa, D. (2020). Lucid: A Practical, Lightweight Deep Learning Solution for DDoS Attack Detection. *IEEE Transactions on Network and Service Management*, 17(2), 876-889. <https://doi.org/10.1109/TNSM.2020.2971776>
- Fadlil, A., Riadi, I., & Aji, S. (2017). Review of detection DDoS attack detection us-ing Naïve Bayes classifier for network forensics. *Bulletin of Electrical Engineering and Informatics*, 6, 140-148. <https://www.unb.ca/cic/datasets/ddos-2019.html>
- University of New Brunswick. *DDoS Evaluation Dataset (CIC-DDoS2019)*. (2021). Retrieved December 20, 2021, from <https://www.unb.ca/cic/da-tasets/ddos-2019.html>
- Setitra, M. A., Fan, M., Agbley, B. L., & Bensalem, Z. E. (2023). Optimized MLP-CNN model to enhance detecting DDoS attacks in SDN environment. *Network*, 3(4), 538-562.
- Wang, C., Zheng, J., & Li, X. (2017). *Research on DDoS attacks detection based on RDF-SVM*. In *Proceedings of the 10th International Conference on Intelli-gent Computation Technology and Automation* (pp.9-12). Changsha, China.
- Zahid Hasan, Md., Zubair Hasan, K. M., & Sattar, Abdus (2018). Burst header packet flood detection in optical burst switching network using deep learning model. *Procedia Computer Science*, 143, 970-977.
- Zhu, M., Ye, K., & Xu, C. Z. (2018). *Network anomaly detection and identification based on deep learning methods*. In M. Luo, & L. J. Zhang (Eds.), *Cloud Computing – CLOUD 2018*. CLOUD 2018. Lecture Notes in Com-puter Science. Cham: Springer.

LONG SHORT-TERM MEMORY DEEP LEARNING MODEL DETECTING DDOS ATTACKS

Phạm Trọng Huynh

Ho Chi Minh City University of Natural Resources and Environment, Vietnam

Corresponding author: Nguyen Duy Vy – Email: nguyenduyvy@gmail.com

Received: June 13, 2024; Revised: December 11, 2024; Accepted: January 23, 2025

ABSTRACT

Recently, Distributed Denial of Service (DDoS) attack threats have become increasingly sophisticated, posing challenges to conventional defense systems. Early detection of attack signs is crucial to protect against and counter these attack threats. The research proposes using a model based on the Long Short-Term Memory (LSTM) deep learning technique to identify DDoS threats in network data packets. This LSTM technique includes selected algorithms and feature extraction, automatically updated during training. Even with a small amount of data, LSTM operates quickly and accurately. The study conducted experiments on the CICDDoS2019 dataset, with results showing the model achieving the following performance metrics: Accuracy reaching 93%, Precision at 96%, Recall reaching 93%, and F1 Score at 94%. The research aims to provide a model capable of processing sequential data and retaining long-term learned information. Integrating the model into network monitoring and security systems can enhance the ability to detect and respond to increasingly complex network attack threats.

Keywords: DDoS; Deep Learning; DoS; LSTM; Machine Learning