

네트워크 플로우 데이터 기반 이상징후 탐지 인공지능 모델 성능 비교

박태정*, 최태정* , 서성관**, 손배훈**, 윤주범***

xowjdpk@gmail.com taeda593@gmail.com *

ths3847@gmail.com sbh960701@gmail.com** jbyun@sejong.ac.kr***

Performance Comparison of Anomaly Detection Using Artificial Intelligence Model based on Network Flow Data

Tae Jeong Park, Tae Jeong Choi* Seong Kwan Seo, Bae Hun Son **

Joo Beom Yun***

Sejong Univ(Undergraduate Student)* Sejong Univ(Graduate Student)**

Sejong Univ(Professor) ***

요약

네트워크 기반 공격의 위험성은 서비스에 치명적이며 사회적, 경제적 손실을 야기한다. 따라서 네트워크를 위협하는 여러 공격에 대한 사전 관리와 예방이 손실을 줄일 수 있는 방법이다. 본 논문에서는 네트워크 플로우 기반 데이터로 이상징후를 탐지하는 인공지능 모델 SVM, AutoEncoder, LSTM AutoEncoder, Isolation Forest를 설계하고 그 성능을 비교하였다. 데이터는 플로우 기반 네트워크 데이터인 CICIDS 2017을 사용하여 인공지능 모델을 학습시키고 테스트했다. SVM, AutoEncoder, LSTM AutoEncoder, Isolation Forest 중 LSTM AutoEncoder 가 F1-Score 95.21 %로 가장 뛰어난 성능을 보였다.

I. 서론

네트워크 트래픽 이상탐지는 서비스 및 서버 보호를 위한 정책 중 가장 중요 한 구성 요소로서 정상 동작을 따르지 않는 플로우를 빠르고 정확하게 탐지하는 것이 무엇보다 중요하다. 불규칙한 패턴을 가진 플로우 데이터를 보다 잘 분류하기 위해서는 정보의 특성을 모델링할 수 있는 신경망 알고리즘이 필요하다.

이에 본 논문에서는 지도학습 모델인 SVM 과 비지도 학습 모델 Autoencoder (AE),LSTM AE(Autoencoder), Isolation Forest(I-Forest) 를 활용 한 네트워크 트래픽 이상탐지 모델을 설계하고 알고리즘 간 성능을 비교하였다.

II. 본론

2.1 데이터 처리

인공지능 모델을 학습시키기 위해서는 많은 양의 데이터 샘플이 필요하다. 실험 한 네트워크 트래픽 이상탐지 신경망을 훈련하고 검증하기 위해 정상데이터와 DoS Golden Eye, DoS Hulk, DoS Slowloris, DoS SlowHTTP Test, Heart bleed로 진단된 데이터 Intrusion Detection Evaluation Dataset (CIC-IDS2017) 을 사용하였다. 인

공지능 모델에 입력으로 사용하기 위한 데이터 전처리 단계는 총 세단계로 구성 하였다.

첫번째 단계는 One-hot-encoding 으로 정상, 이상 데이터 라벨에 대응하는 값을 0 과 1로 구분 하였다. 두번째는 네트워크 플로우 기반 데이터를 Correlation 에 기반하여 Feature 를 추출하였다.

$$x_{i_new} = \frac{x_i - \text{mean}(x)}{\text{stdev}(x)} \quad (1)$$

세번째로 추출 된 Feature 에서 데이터의 값을 식 (1)을 이용하여 데이터의 평균을 0, 분산을 1로 변환하여 표준화하는 과정을 거쳤다.

2.2 실험 인공지능 모델 구조

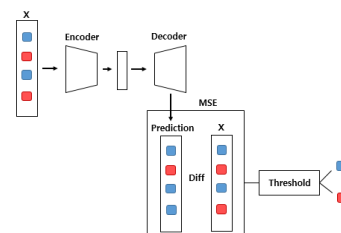


그림 1. AE 와 LSTM AE 모델의 신경망 구조

AE 와 LSTM AE 모델의 신경망은 인코더

계층과 디코더 계층으로 구성되었다. 정상 데이터가 신경망에 입력되면 인코더 층에서 압축되고 디코더 층에서 입력과 똑같이 복원되는 것을 목표로 복원 데이터의 손실을 최소화하는 과정을 통해 정상 데이터의 특징을 학습한다. 그림 1은 테스트 데이터 X가 AE와 LSTM AE 모델의 학습된 신경망에 입력되는 과정을 나타낸 것이다. X는 정상과 비정상 데이터로 구성되어 있다. X가 신경망을 거쳐 복원값(Prediction)이 반환되면 원본 X와 복원값을 비교하여 그 차이가 임계값(Threshold) 보다 클 경우 비정상 데이터로 분류하고, 작은 경우 정상 데이터로 분류하며 이상탐지를 수행한다.

SVM 모델은 정상과 비정상 데이터를 입력받아 결정 경계를 정의하고 테스트 데이터 입력시 결정 경계와 비교하여 정상과 비정상으로 분류하여 이상탐지를 수행한다.

Isolation Forest(I-Forest)는 랜덤하게 차원을 선택하여 공간을 분할한다. 분할된 공간에 군집화 된 정상치는 공간내에서 많은 공간분할이 필요하지만 이상치는 정상 군집에서 멀리 떨어져 있으므로 적은 횟수의 공간 분할만으로 이상치를 고립 (isolation)시킬 수 있다. 이 과정은 그림2의 Decision Tree로 표현되는데 정상치 일수록 완전히 고립시키도록 Tree의 깊이가 높아진다.

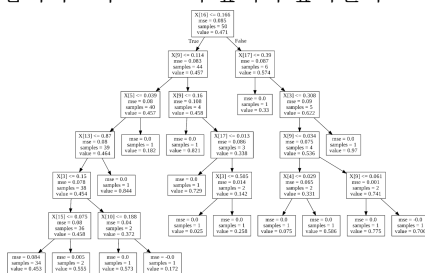


그림2. Isolation Forest Decision Tree

2.3 실험 및 성능비교

실험 데이터는 학습 60%, 검증 10%, 테스트 30%로 분할하였으며 검증의 타당성을 높이기 위해 모든 신경망 알고리즘에 대해 10겹 교차검증을 수행하였다. 표1은 10겹 교차검증을 수행한 알고리즘 이상탐지 성능 평균을 나타낸다. 이상 탐지 신경망의 성능 평가 지표로 Accuracy, Precision, Recall, F1 Score를 사용하였다.

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
AE	92.83	100	89.71	94.53
I- Forest	63.52	70.86	82.40	76.19
SVM	84.57	75.91	96.46	84.96
LSTMAE	95.33	95.41	95.22	95.31

표 1. 알고리즘 이상탐지 성능

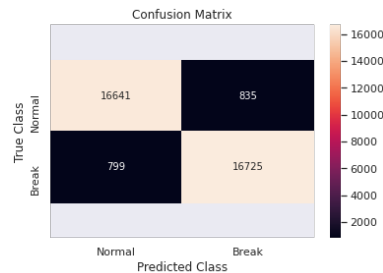


그림 3. LSTM-AE 이상 탐지 혼동행렬

LSTM AE를 적용한 이상탐지에서 95.33% Accuracy를 기록하였고 이는 다른 알고리즘과 비교하여 2% 이상 높은 수치이다. LSTM AE는 Recall과 F1-score에서도 다른 알고리즘의 성능보다 높은 수치를 보여준다. 그림2는 LSTM-AE로 분류한 CICIDS 2017 데이터셋의 혼동행렬을 나타낸다. 약 5% 정도의 오분류 확인할 수 있다.

III. 결론

본 논문에서는 네트워크 플로우 데이터를 사용하여 비정상 트래픽을 탐지하기 위한 네 가지 이상탐지 인공지능 모델을 비교했다. SVM, AE, LSTM AE, Isolation Forest 모델을 학습시켜 네트워크 플로우 데이터에 대한 이상탐지 성능을 측정하였다.

실험 결과 LSTM AE가 다른 인공지능 알고리즘보다 높은 성능을 보여준다. 이는 네트워크 이상탐지 알고리즘에 LSTM AE가 적용될 수 있음을 의미한다. 이후 LSTM AE 성능개선/최적화를 위한 연구를 진행할 계획이다.

ACKNOWLEDGMENT

본 연구는 과학기술정보통신부 및 정보통신기획평가원의 대학ICT연구센터지원사업의 연구 결과로 수행되었음 (IITP-2021-2018-0-01423)

참고 문헌

- [1] K.J. Kim, "A survey of deep learning-based network anomaly detection," Cluster Computing, vol.27, pp. 949-961, Jan. 2019.
- [2] T. N. Sainath, O. Vinyals, A. Senior and H. Sak, "Convolutional, long short-term memory, fully connected deep neural networks," Acoustics, Speech and Signal Processing, IEEE, pp. 4580-4584, 2015.
- [3] Lombardi, M. Milano, and L. Benini, "Anomaly detection using autoencoders in high performance computing systems", In Proceedings of the AAAI Conference on Artificial Intelligence, vol. 33, 2019.