

トラフィックデータのエントロピーに基づく アノマリ型 IDS の構築

鷺坂 典雅^{1,a)} 青木 茂樹^{1,b)} 宮本 貴朗¹

概要：近年，組織や個人など特定のターゲットを狙って行われる標的型攻撃の増加や巧妙化が大きな脅威となっており，効果的な検知技術の実現が求められている．標的型攻撃には未知の攻撃が含まれていることが多く，未知の攻撃の検知に有効であるアノマリ型 IDS (Intrusion Detection System) の重要性が高まっている．本研究では，日常業務における組織内の通信には固有のパターンが存在すると仮定し，通常時の通信を学習して，不審な通信を検出する手法を提案する．まず，通常時の通信を単位パケット数ごとに分割し，分割した各区間ごとにパケットのヘッダから IP アドレス，ポート番号等の特徴量を抽出する．抽出した各特徴量からエントロピーを算出し，特徴ベクトルを生成する．次に，生成した特徴ベクトルを学習データとしてクラスタリングする．その後，テストデータから同様に特徴ベクトルを生成し，学習データのクラスタに所属するか否かを判別する．クラスタに所属しない場合，学習データとは類似しない通信が行われたと判断し，異常として検出する．実験では，MWS データセットを用いて本手法の有効性を確認した．

キーワード：異常検出，クラスタリング，IDS，エントロピー

Construction of anomaly based IDS based on entropy of traffic data

NORIMASA SAGISAKA^{1,a)} SHIGEKI AOKI^{1,b)} TAKAO MIYAMOTO¹

Abstract: Increase in APT (Advanced Persistent Threat) is a serious threat and requires effective detection technology. Since APT include unknown attacks, anomaly-based IDS (Intrusion Detection System) is considered effective in detecting them. In this paper, we propose a method to detect suspicious traffic by learning normal traffic in organization. First, traffic data is divided by number of packets, and features such as IP addresses are extracted from the packets for each divided section. We calculate entropy from each extracted features and generate feature vectors. We cluster the generated feature vectors. Next, we generate a feature vector from new test data and select the closest cluster. We detect anomaly when distance between the feature vector and the cluster center is greater than threshold. In the experiment, we confirmed effectiveness of our method using MWS dataset.

Keywords: Anomaly Detection, Clustering, IDS, Entropy

1. はじめに

近年，組織や個人など特定のターゲットを狙って行われる標的型攻撃の増加や巧妙化が大きな脅威となっており，

効果的な検知技術の実現が求められている．標的型攻撃では未知の攻撃が利用されることが多く，実環境でよく用いられている Snort [1] や Suricata[2], Zeek[3] などのシグネチャ型 IDS (Intrusion Detection System) では，予め攻撃のパターンをシグネチャとして用意しておく必要があるために，検知できないことが問題となっていた．そこで，未知の攻撃の検知に有効であるアノマリ型 IDS の重要性が高まっている．

¹ 大阪府立大学大学院人間社会システム科学研究科
Graduate School of Humanities and Sustainable System Sciences, Osaka Prefecture University

a) saa01119@edu.osakafu-u.ac.jp

b) aoki@kis.osakafu-u.ac.jp

アノマリ型 IDS は一般的にシグネチャ型 IDS より誤検知は増えるものの、シグネチャ型 IDS では検知できない未知の攻撃を検知できる。代表的なアノマリ型 IDS に関する研究 [4] では、正常なトラフィックを学習し、学習した正常状態からの外れ値を異常として検知している。文献 [5], [6] では、トラフィックをパターンごとにクラスタリングして学習し、クラスタに属さないトラフィックを異常として検出している。これらの手法では、正常のトラフィックを特徴量の類似するトラフィックごとにクラスタリングし、複数の状態として学習しているため、単一の状態として扱う手法よりも高精度に異常を検知できる。更に、文献 [7], [8], [9] ではパケットのヘッダ情報から抽出した特徴量のエントロピーに注目し、未知の異常を検出する手法を提案している。

本研究では、組織内における日常業務の通信には固有のパターンが存在すると仮定し、通常時の通信を学習して、不審な通信を検出する手法を提案する。まず、通常時の通信を単位パケット数ごとに分割し、分割した各区間ごとにパケットのヘッダから IP アドレス、ポート番号等の 11 種類の特徴量を抽出する。抽出した各特徴量のエントロピーを算出し、特徴ベクトルを生成する。次に、生成した特徴ベクトルをクラスタリングして学習する。その後、テストデータから同様に特徴ベクトルを生成し、学習したクラスタに所属するか否かを判別する。クラスタに所属しない場合、学習データとは類似しない通信が行われたと判断し、異常として検出する。実験では、MWS データセットの BOS データセット [11] に本手法を適用して有効性を確認した。

以下、2 節で、関連研究について述べ、3 節では特徴の抽出方法とエントロピーの算出方法、クラスタリングおよび異常検知手法について説明する。4 節では、実験結果と考察について述べ、5 節でまとめと今後の課題について述べる。

2. 関連研究

本研究に関連する従来研究として、クラスタリング結果に基づいて異常を検出するアノマリ型 IDS の手法である文献 [5][6] と、特徴量にエントロピーを用いて異常を検知する手法である文献 [7], [8], [9] について述べる。

文献 [5] では、ネットワークトラフィックの正常状態は単一ではなく、トラフィックの特徴ごとに複数の状態に分類できると考えて、正常状態を複数定義し、各状態との違いを基に異常を検出する手法を提案している。この手法では、異常を含まないデータから単位時間あたりの ICMP や TCP パケット数等を計測して特徴ベクトルを生成し、生成した特徴ベクトルをクラスタリングする。メンバが少ないクラスタは削除し全てのクラスタにおいて閾値以上のメンバ数となるまでクラスタリングを繰り返す。クラスタリング結果を正常状態として定義し、新たに観測されたデー

タから同様に特徴を抽出し、正常クラスタとの距離を基に異常を識別している。

文献 [6] の手法では、パケットのヘッダ情報から 61 次元の特徴量を抽出し、主成分分析法によって特徴の次元を圧縮した後、クラスタリングしている。更に、各クラスタにシグネチャ型 IDS による攻撃の検出結果をラベル付けすることにより、異常を検出するだけではなく、異常の種類を通知可能なアノマリ型 IDS を構築している。

文献 [5], [6] の手法では、抽出した特徴量をそのまま用いて異常検出しているために、DoS 攻撃や DDoS 攻撃などの様にパケット数などに特徴が現れる異常については安定して検出できるものの、本研究で対象としている標的型攻撃で用いられる、マルウェアと C2 サーバ間のビーコン通信などパケット数の急激な増加を伴わない異常の検出は難しいと考えられる。

一方文献 [7] では、パケットのエントロピーに基づく異常検出手法が提案されている。この手法ではまず、到着順に並べたパケットを単位パケット数ごとに区間に分割する。各区間で IP アドレスやポート番号などの特徴量を抽出し、それぞれの特徴量の発生確率を求め、求めた発生確率からエントロピーを算出する。その後、エントロピーの時系列変化に着目した EMMM 法 (Entropy based Multidimensional Mahalanobis distance Method) により、エントロピーが大きく変化する時間を攻撃などが含まれている異常区間として検出している。

文献 [8] では文献 [7] の手法と同様に、パケットのヘッダから抽出した特徴からエントロピーを算出して特徴ベクトルを生成し、生成した特徴ベクトルに対して逐次学習を行う外れ値検出モデルと、SOM を組み合わせることで、攻撃の変化に対応した持続的な異常検出手法を提案している。

文献 [9] では、標的型攻撃の侵入初期段階に用いられる遠隔操作ツールである RAT (Remote Access Tool) を検出する手法を提案している。この手法では、RAT が確立した C2 サーバとの通信トラフィックのパケットは、送信間隔に偏りがあることに注目し、パケット到着時間間隔の偏りをエントロピーにより数値化し、特徴量とすることで RAT と正常なアプリケーションを識別する手法を提案している。

文献 [7], [8], [9] の手法では、抽出した特徴量のエントロピーを用いることによって、特徴量の値自体の変化ではなく特徴量の発生確率の変化に注目している。そのため、パケット数などが急激に増加しない、マルウェアと C2 サーバとのビーコン通信などについても的確に特徴を捉えることができ、異常として検出できると考えられる。

3. 提案手法

本稿では、パケットのヘッダから抽出した特徴量を基に算出したエントロピーを特徴ベクトルとし、特徴ベクトル

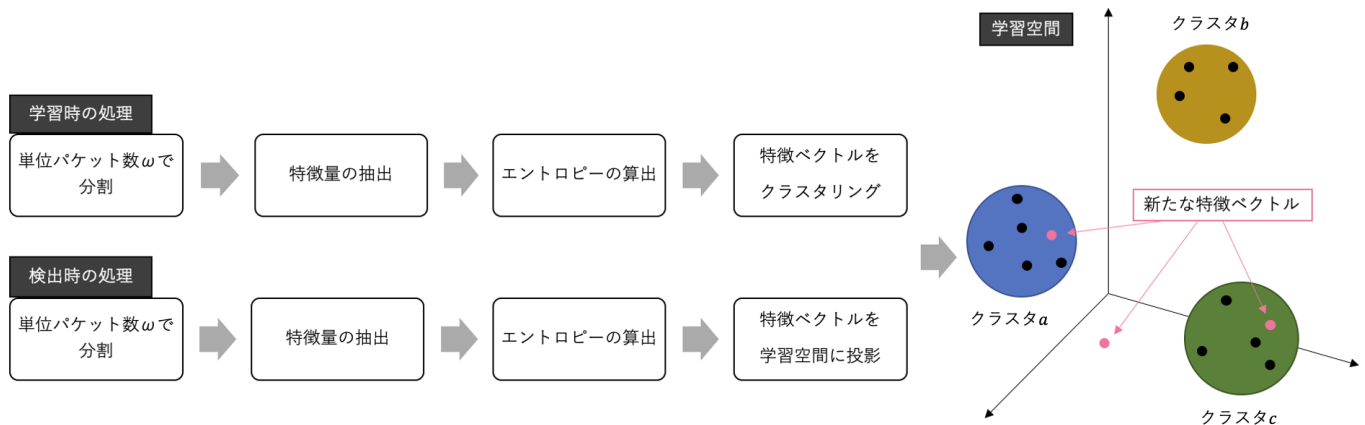


図 1 提案手法の概要

Fig. 1 Outline of Proposed Method.

をクラスタリングすることによって、異常を検知する手法を提案する。

本手法による異常検出の流れを図 1 に基づいて説明する。本手法は学習処理と検出処理に分かれている。学習処理ではまず、学習用トラフィックデータを単位パケット数ごとに分割して区間とし、各区間に含まれるパケットのヘッダから 11 種類の特徴量を抽出する。次に、抽出した特徴量のエントロピーをそれぞれ算出して特徴ベクトルとする。その後、特徴ベクトルをクラスタリングすることで、トラフィックを複数のパターンに分類して学習する。検出処理では、学習処理と同様にトラフィックデータから区間ごとにエントロピーに基づく特徴ベクトルを生成し、生成した特徴ベクトルを学習した空間に投影する。投影した特徴ベクトルが、学習したどのクラスタにも所属しない場合、その区間の通信は学習用トラフィックデータと類似しないと判断し、異常として検出する。

3.1 特徴ベクトルの抽出

トラフィックデータから特徴ベクトルを抽出する。特徴ベクトルを抽出する方法として、セッションごとに抽出する方法と一定の区間ごとに抽出する方法が考えられる。セッションごとに特徴量を抽出する場合、どの通信が攻撃であるかを容易に特定することができるが、トラフィック量の多いネットワークではセッション数が非常に多くなるために、計算コストが膨大になる問題がある。トラフィックデータを区間ごとに分割して特徴量を抽出する方法では、攻撃通信を特定することは難しいものの、セッションごとに抽出する場合よりも計算コストを大幅に低減するこ

とができるため、トラフィック量が多いネットワークへの適用が容易である。そこで、本手法では区間ごとに特徴を抽出することとする。

区間ごとに特徴を抽出する手法としては、文献 [6] のように単位時間でトラフィックデータを分割する手法と、文献 [8] のように単位パケット数で分割する手法が提案されている。単位時間で分割する場合、攻撃が行われた時間の特定が容易であったり、DDoS 攻撃や DoS 攻撃などの様に特徴がパケット数に現れる攻撃を検出しやすいなどの利点があるものの、特徴ベクトルの生成にエントロピーなどを適用しにくいと、本研究で対象としている、標的型攻撃のビーコン通信などのようにパケット数の急激な変化を伴わない異常の検出が難しいことが課題となっていた。本手法では、単位パケット数でネットワークトラフィックを区間に分割し、区間から抽出した特徴量のエントロピーを特徴ベクトルとして利用する。単位パケット数で区間に分割することにより区間ごとのパケット数が一定となるためエントロピーを適用しやすく、特徴量の発生確率の変化などに現れる特徴を扱うことができるため、標的型攻撃のビーコン通信などを検出できる。また、DoS 攻撃や DDoS 攻撃などの様にパケット数に特徴が現れる攻撃についても、特徴量の発生確率が変化するため、異常として検出できると考えられる。

あるネットワークに対する攻撃を検知するために、対象とするネットワークから収集したトラフィックデータから特徴量を抽出し、特徴ベクトル I_t を生成する。特徴ベクトル I_t の生成手順を以下に示す。

- (1) パケットを到着時間順に並べ、連続するパケット列を

表 1 抽出する特徴量一覧
Table 1 List of Features.

送信元 IP アドレス
宛先 IP アドレス
送信元ポート番号
宛先ポート番号
パケットバイト数
プロトコル
フラグメント ID
フラグメント用フラグの状態
TTL 値
パケット到達間隔
コントロールフラグ

単位パケット数 ω で分割する.

- (2) 分割した各区間のパケットのヘッダから, 表 1 に示す IP アドレスやポート番号等 11 種類の特徴量を取得する.
- (3) それぞれの特徴量で出現したシンボルごとに出現回数をカウントする. シンボル i の出現回数 x_i から出現確率 $P_i = x_i/\omega$ を算出する. ポート番号で出現確率を算出する例を表 2 に示す. 表 2 の例では, 100 パケットの区間のうち 1 つ目のシンボルである 22 番ポートが 15 回出現している. そのため出現確率 $P_1 = 15/100$ となる
- (4) シンボルの総数を m とし, 式 (1) によりエントロピー h を求める.

$$h = - \sum_{i=1}^m P_i \log P_i \quad (1)$$

表 2 に示す例でのエントロピーは式 (2) のように算出する.

$$h = -\left\{\left(\frac{15}{100} \log \frac{15}{100}\right) + \left(\frac{20}{100} \log \frac{20}{100}\right) + \left(\frac{50}{100} \log \frac{50}{100}\right) + \left(\frac{10}{100} \log \frac{10}{100}\right) + \left(\frac{5}{100} \log \frac{5}{100}\right)\right\} = 0.5789 \quad (2)$$

- (5) 算出した 11 次元のエントロピーを座標値とする特徴ベクトル $\mathbf{I}_t = (h_1, h_2, \dots, h_{11})$ を生成する.
- (6) 各区間で上記の手順を繰り返すことで特徴ベクトル $\mathbf{I}_t, \mathbf{I}_{t+1}, \mathbf{I}_{t+2}, \dots$ を抽出する.

表 2 シンボルの出現確率の算出例
Table 2 Calculation example of symbol appearance probability.

シンボル i	出現回数 x_i	出現確率 $P_i = x_i/\omega$
22 番ポート	15	15/100
53 番ポート	20	20/100
80 番ポート	50	50/100
110 番ポート	10	10/100
161 番ポート	5	5/100

3.2 特徴ベクトルのクラスタリング

区間 $t, t+1$ に含まれる通信の内容が類似する場合, 出現する特徴量の発生確率が近くなり, エントロピーが類似するため, 抽出した特徴ベクトル $\mathbf{I}_t, \mathbf{I}_{t+1}$ 間の距離は小さくなる. 異なる内容の通信が含まれる場合, 出現する特徴量の発生確率が変化するためエントロピーが異なる値となり, 特徴ベクトル間の距離は大きくなる. そこで, 分割した各区間の特徴ベクトルを投影する 11 次元の空間を H 空間とし, H 空間上での座標値 $\mathbf{I}_t = (h_1, h_2, \dots, h_{11})$ を基にクラスタリングすることで, 含まれる通信の内容ごとに区間を分類する.

本研究では Mean-Shift 法を用いてクラスタリングする. クラスタリングでよく用いられる k -means 法のようなクラスタ解析手法では, クラスタ数が事前にわかっていることを前提としているが, Mean-Shift 法はクラスタ数を事前知っておく必要がない. そのため, 学習データの通信パターン種類の数があらかじめわからない今回のような場合に適した手法である.

3.3 異常検出

新たに観測されたトラフィックデータを, 学習時と同様に単位パケット数 ω で分割し, それぞれの区間から特徴量を抽出して, エントロピーを算出し, 特徴ベクトルを生成する. 生成した特徴ベクトルを H 空間に投影する. 区間に学習時と同様の内容の通信のみが含まれる場合, 特徴ベクトルは学習したクラスタに投影される. 図 2 に示すようにクラスタ a , クラスタ b が存在するとき, 新たな特徴ベクトル $\mathbf{I}_t$ が各クラスタの重心との距離において, クラスタ a の重心との距離が最も近いと算出されたとする. 特徴ベクトル $\mathbf{I}_t$ とクラスタ a の重心との距離を e とする. また, クラスタ a 内の特徴ベクトル $\mathbf{I}_a, \mathbf{I}_{a+1}, \mathbf{I}_{a+2}, \dots$ の中で, 重心との距離が最も大きい特徴ベクトル $\mathbf{I}_{a+2}$ と重心との距離を f とする. $e \leq f$ が成り立つため, 特徴ベクトル $\mathbf{I}_t$ はクラスタ a に属すると判断する.

次に, 新たな特徴ベクトル $\mathbf{I}_{t+1}$ が各クラスタの重心との距離において, クラスタ b の重心との距離が最も近いと算出されたとする. 特徴ベクトル $\mathbf{I}_{t+1}$ とクラスタ b の重心との距離を e' とする. また, クラスタ b 内の特徴ベクトル

表 3 進行度の説明

Table 3 Explanation of Progress.

進行度	説明
1, 2	通信発生なし
3, 4, 5	通信は発生したが, C&C サーバとの攻撃通信は不成立
6, 7, 8	通信発生かつ, C&C サーバとの攻撃通信が成立

C2 サーバとの通信が区間内に 15% 以上含まれる区間を異常区間と定義した. 実験データには, マルウェアは実行したが通信は発生しなかった進行度 2 の 2017 年 8 月 17 日に収集された 24 時間のトラフィックデータと, マルウェアを実行し C2 サーバとの通信が継続的に観測できた進行度 8 の 2018 年 1 月 23 日と 2018 年 1 月 24 日にそれぞれ収集された 24 時間のトラフィックデータを用いた.

4.2 クラスタリング実験

本手法で正常区間と異常区間进行分类できることを確認するための実験を行った. 実験では 2018 年 1 月 23 日に収集された進行度 8 のトラフィックデータを用いて実験を行った. クラスタリングにより分類された各クラスターのうち, 異常区間が半数以上を占めるクラスターを異常クラスター, 半数未満のクラスターを正常クラスターとした. なお, この実験では特徴抽出の際の単位パケット数を $\omega = 100$ としている. クラスタリング実験の評価指標には F 値と特異度を用いる. F 値は実際に正常区間であるもののうち, 正常クラスターに分類されたものの割合である再現率 (Recall) と, 正常クラスターに分類された区間のうち, 実際に正常区間であるものの割合である適合率 (Precision) の調和平均である. 式 (3) に F 値の算出式を示す.

$$\text{F-measure} = \frac{2 * \text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall}} \quad (3)$$

特異度は, 実際に異常区間であるもののうち, 異常クラスターに分類されたものの割合である.

Mean-Shift でクラスタリングするときのパラメータである, カーネル幅 σ の値を変化させた際のクラスタリング実験結果を図 3, 図 4 に示す. 図 3 は分類精度を確認するために F 値と特異度を調べた結果であり, 横軸にカーネル幅 σ , 縦軸に F 値と特異度を示している. また, 図 4 は σ を変化させたときのクラスター数を調べた結果を示しており, 横軸にカーネル幅 σ , 縦軸にクラスター数を示している. 図 3 より, F 値は σ によらず高い値を示したが, 特異度は σ が 0.2 を超えると低下することを確認した. F 値が常に高い値を示しているのは, 実験に用いたデータに含まれている異常区間が正常区間と比べて少ないために, 正常の影響を大きく受けたことが原因だと考えられる. また, 特異度が $\sigma > 0.2$ のときに低下するのは, σ が大きくなるほど各クラスターの範囲が大きくなり, 正常区間と異常区間が混

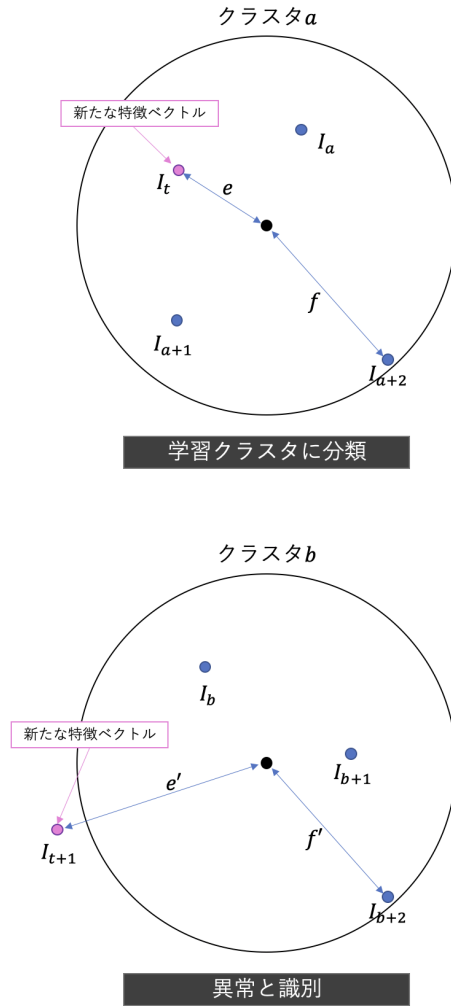


図 2 異常検出の例

Fig. 2 Example of Anomaly Detection.

ル $I_b, I_{b+1}, I_{b+2}, \dots$ の中で, 重心との距離が最も大きい特徴ベクトル I_{b+2} と重心との距離を f' とすると, $e' > f'$ が成り立つため, 特徴ベクトル I_{t+1} はどのクラスターにも属しないと判断し, 異常と識別する.

4. 実験

4.1 実験条件

本手法の有効性を確認するために実験を行った. 実験では, 標的型攻撃を含むデータセットである MWS2018 データセット [11] 中の BOS データセットを用いて, クラスタリング及び異常検出実験を行なった. BOS データセットは組織内ネットワークへの侵害活動を想定し, 1 から 8 までの進行度が定義された動的活動観測トラフィックデータのデータセットである. 表 3 に進行度の説明を示す. 本実験ではトラフィックデータに含まれる C2 サーバとの通信を異常と定義し, それ以外の通信を正常と定義した. また,

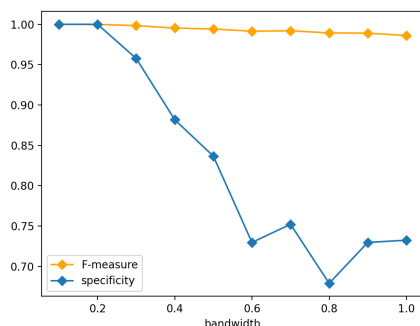


図 3 F 値と特異度

Fig. 3 F-measure and specificity

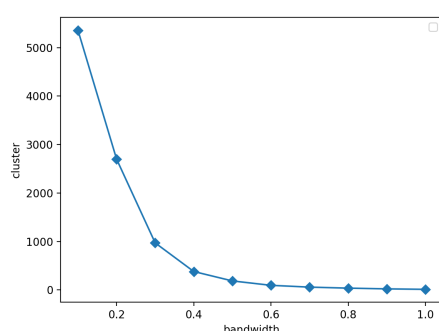


図 4 クラスタ数

Fig. 4 Number of cluster

在するクラスタが増加するためだと考えられる。以上の実験結果から、本手法を用いて正常区間と異常区間を分類できることを確認できた。

図 4 より σ の値が 0.3 を下回るとクラスタ数が急激に増加している。これは、カーネル幅 σ を小さくするほど各クラスタの範囲が狭くなり、細かな分類となるためである。クラスタのメンバ数を確認したところ、 σ が小さい場合には各クラスタに含まれる区間数が減少し、単独区間からなるクラスタが増加していることを確認した。 σ に小さすぎる値を設定した場合、3.3 節の手法によりに異常を検出する際に、正常通信を異常として誤検知しやすくなる可能性がある。

4.3 検出実験

本手法で標的型攻撃を検出できることを確認するために異常検出実験を行った。異常検出実験では、学習データとして 2017 年 8 月 17 日に収集された進行度 2 のトラフィックデータを用いた。テストデータには 2018 年 1 月 23 日と 2018 年 1 月 24 日に収集された進行度 8 のトラフィックデータを用いた。本手法と既存手法 [6] を、クラスタリングの際のカーネル幅 σ を変化させて描画した ROC 曲線と AUC 値により比較した。ここで、ROC 曲線下面積である AUC 値は 0 から 1 までの値を取り、1 に近づく

ほど優れた識別器であることを示す。文献 [6] の手法では特徴抽出の区間を単位時間で分割しているが、文献 [6] の手法で分割方法を単位パケット数 ω とした場合の結果も併せて比較する。なお、本実験での単位パケット数 ω は 100 としている。2018 年 1 月 24 日のテストデータでの本手法における検出結果と、従来手法における検出実験結果を図 5 に示す。実験の結果、本手法は単位時間ごとに特徴量を抽出して異常を検出する文献 [6] の手法や文献 [6] の手法の特徴抽出方法を、単位パケット数ごとに抽出するように変更した場合よりも高い AUC 値が得られたため、単位パケット数ごとに特徴量を抽出し、抽出した特徴量のエントロピーに基づいて異常を検出する手法が標的型攻撃の検出に有効であることを確認できた。

次に、本手法の区間分割の単位パケット数 ω を 50, 100, 200, 300 と変化させたときの異常検出精度を、ROC 曲線および AUC 値で確認する実験を行った。学習データには 2017 年 8 月 17 日に収集された進行度 2 のトラフィックデータを用い、テストデータとして 2018 年 1 月 23 日に収集されたトラフィックデータを用いた際の検出実験結果を図 6 に示し、2018 年 1 月 24 日に収集したトラフィックデータを用いた際の検出実験結果を図 7 に示している。検出実験結果より、用いるテストデータにより精度に差はあるものの、どちらのテストデータを用いても単位パケット数 ω を大きく設定した方が識別精度が向上することを確認した。 ω の値が小さい場合、サイズの小さなファイルを Web サーバに複数回アップロードされたときの特徴ベクトルと、HTTP に対する DoS 攻撃などを受けた場合の特徴ベクトルが類似してしまうため、正常と異常を正しく識別できずに識別精度が低下してしまうが、 ω の値を 300 程度に設定した場合には、このような問題が発生しないために識別精度が向上したと考えられる。

図 6, 図 7 に示す結果から、 $\omega = 50$ から $\omega = 100$ に変化させたときと比較して、 $\omega = 200$ から $\omega = 300$ に変化させたときの AUC 値の増加量は少ないことが確認できる。そのため、 $\omega = 300$ 程度より大きな値を設定したとしても、識別精度の大幅な向上は見込めないと考えられる。また、単位パケット数に大きすぎる値を設定した場合、標的型攻撃のビーコン通信のようにパケット数の少ない攻撃の特徴が特徴ベクトルに十分に反映されず、異常検出精度が低下する可能性が考えられる。今後の課題として、最適な単位パケット数 ω の設定方法について検討することなどが挙げられる。

5. まとめ

本稿では、単位パケット数 ω ごとに分割したパケットのヘッダ情報から抽出した 11 次元の特徴量からエントロピーを算出し、特徴ベクトルを生成した。学習データから生成した特徴ベクトルを基にクラスタリングを行い、新た

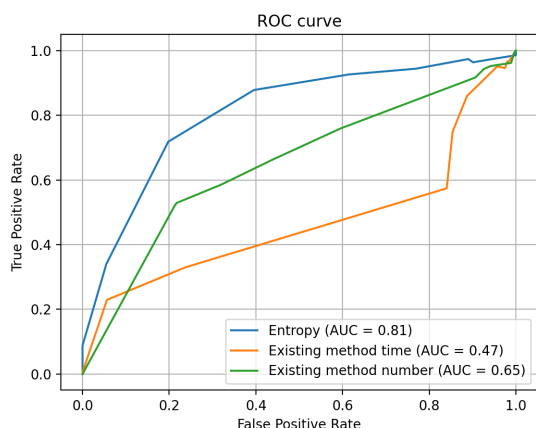


図 5 本手法と従来手法の検出結果の比較 ($s = 15\%$)

Fig. 5 Comparison between results of Proposed Method and Conventional Method. ($s = 15\%$)

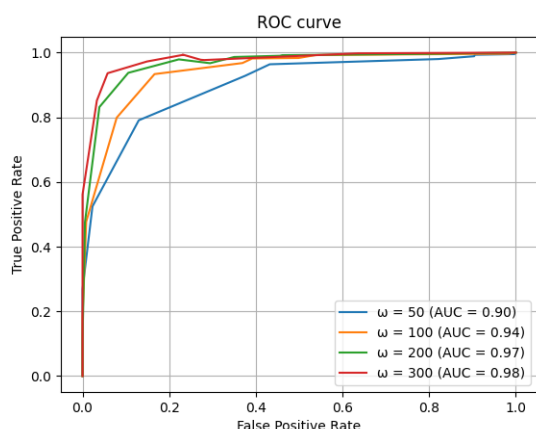


図 6 ω の値ごとの比較結果 (2018/1/23 $s = 15\%$)

Fig. 6 Comparison results when ω is changed. (2018/1/23 $s = 15\%$)

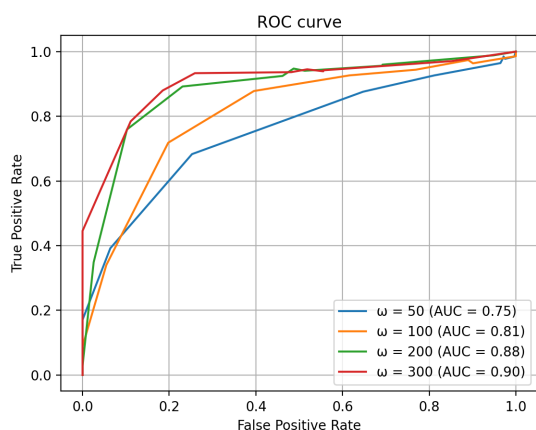


図 7 ω の値ごとの比較結果 (2018/1/24 $s = 15\%$)

Fig. 7 Comparison results when ω is changed. (2018/1/24 $s = 15\%$)

なトラフィックデータから抽出した特徴ベクトルを学習した空間に投影することで未知の異常通信を検知するアノマリ型 IDS を構築した。

実験では標的型攻撃の通信を含む MWS データセット中の BOS データセットを用いて本手法の有効性を確認した。実験の結果、標的型攻撃の通信を高い AUC 値で異常として検出できることを確認した。また、従来手法 [6] との比較実験により、単位パケット数でパケットを分割しエントロピーを算出して特徴ベクトルとする本手法の有効性を確認できた。

今後の課題としては、クラスタリングの際のパラメータであるカーネル幅 σ や単位パケット数 ω の設定方法の検討や、実運用中のトラフィックに適用した際の検出精度の検証などが挙げられる。

参考文献

- [1] Snort, <<https://www.snort.org/>>(参照 2020-08-11).
- [2] Suricata, <<http://suricata-ids.org/>>(参照 2020-08-11).
- [3] Zeek, <<https://zeek.org/>>(参照 2020-08-11).
- [4] 佐藤陽平, 和泉勇治, 根元義章: 複数の検出モジュールによるネットワーク異常検出の高精度化, 信学技報, NS2004-144, pp.45-48 (2004).
- [5] 平松尚利, 和泉勇治, 角田裕: 複数の通常状態を用いたネットワーク異常検出, CS2006-32, pp.61-66 (2006).
- [6] 谷澤 俊樹, 青木 茂樹, 宮本 貴朗: パケットのヘッダ情報に注目したアノマリ型 IDS とシグネチャ型 IDS を組み合わせた未知の異常検出, コンピュータセキュリティシンポジウム講演論文集, 3B4-1 (2017).
- [7] 小島 俊輔, 中嶋 卓雄, 末吉 敏則: エントロピーベースのマハラノビス距離による高速な異常検知手法, 情処学論, vol.52, no.2, pp.656-668 (2011).
- [8] 二瓶 凌輔, 青木 茂樹, 宮本 貴朗: 自己組織化マップによるネットワークトラフィックの逐次学習と異常検出, コンピュータセキュリティシンポジウム論文集, pp.1414-1421 (2019).
- [9] 宇野 真純, 石井 将大, 猪俣 敦夫, 新井 イスマイル, 藤川 和利: エントロピーを用いた初期侵入段階における RAT の通信検知手法の考察, 電子情報通信学会論文誌 B, J101-B, pp.220-232 (2018).
- [10] 山田明, 三宅優, 田中俊昭: 亜種攻撃を検知できる侵入検知システム, 電子情報通信学会技術報告, ISEC2004-31, pp.119-126(2004).
- [11] 高田雄太, 他: マルウェア対策のための研究用データセット MWS 2018 Datasets, 情報処理学会, Vol.2018CSEC-82, No.38 (2018).