

クロスファイア攻撃に対して脆弱なエリアの選定法

王 天嶼[†] 上山 憲昭^{††}

[†] 立命館大学大学院 情報理工学研究科

〒 525-8577 滋賀県草津市野路東 1-1-1

^{††} 立命館大学 情報理工学部

〒 525-8577 滋賀県草津市野路東 1-1-1

E-mail: [†]gr0633xx@ed.ritsumei.ac.jp, ^{††}kamiaki@fc.ritsumei.ac.jp

あらまし CFA (Crossfire Attack) は、サーバではなくリンクを標的にするという特徴により、他の DDoS (Distributed Denial of Service) 攻撃と区別される。CFA は検知の難しさから、ネットワークに大きな被害をもたらす可能性がある。しかし、従来のサーバでの攻撃検知に依存する防御方式を用いて CFA を防御することは困難である。CFA に対するより効果的な防御を実現するには、攻撃が開始される前にサーバが存在するネットワークポロジに対する防御を実施することが考えられる。ただし攻撃者が選択するターゲットエリアは予測不可能であるため、CFA 攻撃を事前に予測して防御することは難しい。そこで本稿では、CFA の影響を評価する尺度を定義し、各ネットワークポロジの各エリアを評価することで、CFA 攻撃を受けやすいネットワークポロジの脆弱なエリアを明らかにするアルゴリズムを提案する。提案アルゴリズムにより、CFA 攻撃を受けやすいネットワークポロジの脆弱なエリアを効率的に見つけることができる。提案アルゴリズムの脆弱なエリア選択の精度を評価することで、アルゴリズムの有効性を確認する。

キーワード CFA, ネットワークポロジ, 脆弱性

Identification Method of Vulnerable Target Areas for Crossfire Attacks

Tianyu WANG[†] and Noriaki KAMIYAMA^{††}

[†] Graduate School of Information Science and Engineering, Ritsumei University

1-1-1, Nojihigashi, Kusatsu, Shiga 525-8577

^{††} College of Information Science and Engineering, Ritsumeikan University

1-1-1, Nojihigashi, Kusatsu, Shiga 525-8577

E-mail: [†]gr0633xx@ed.ritsumei.ac.jp, ^{††}kamiaki@fc.ritsumei.ac.jp

Abstract CFA (Crossfire Attack) is distinguished from other DDoS (Distributed Denial of Service) attacks by its characteristic of targeting links rather than servers. CFA can cause significant damage to networks due to their difficulty in detection. However, traditional defense models that rely on detecting DDoS attacks at servers are no longer effective for CFA due to focusing on links. A possible way to provide a more target-oriented defense against CFA is to deploy a defense against the network topology in which the server resides before the attack is launched. However, since the target areas chosen by an attacker are unpredictable, it is difficult to predict CFA attacks in advance and implement a defense plan. Therefore, in this paper, we define a measure to evaluate the impact of CFA and analyze several network topologies to identify the characteristics of vulnerable areas of network topologies susceptible to CFA attacks, and we propose an algorithm to select vulnerable areas using these characteristics. The proposed algorithm can efficiently find vulnerable areas of network topologies susceptible to CFA attacks. Furthermore, we evaluate the accuracy of the proposed algorithm in selecting vulnerable areas to confirm the effectiveness of the proposed algorithm.

Key words CFA, Network topology, Vulnerability

1. はじめに

DDoS (Distributed Denial of Service) 攻撃とは、大量のデータパケットやリクエストを対象サーバに送信し、そのサーバを過負荷にさせてネットワークサービスを利用不能にする行為である。特定のホストに対する DDoS 攻撃に対しては、サーバとネットワークとの接続部分にファイアウォールを設置し攻撃パケットを規制したり、ネットワーク事業者が攻撃を受けている Prefix を BGP ルータで規制するなどの対処法が用いられ

ている [1]。しかし近年、ターゲットエリアに至るネットワークのリンクを高負荷とすることで、ターゲットエリア内のホストを通信不能状態とする Crossfire Attack (CFA) の問題が指摘されている [2]。一般的な DDoS 攻撃とは異なり、CFA の主な特徴は、攻撃の目標がサーバではなく、ネットワーク内のリンクである。目標サーバエリア外のリンクに過負荷をかけることで、ターゲットエリア (TA) と外部との通信を遮断し、TA 内のホストへの通信を外部から遮断してサービスを妨害することを目的とする。

CFA では攻撃ターゲットリンクを選定するために攻撃に先立ち、攻撃に用いる大量のボットとターゲットエリア内の複数のサーバに対し traceroute を行うことで、ターゲットエリアと外部のネットワークとを繋ぐリンクを探索する。このような探索フェーズの後、ボットから送信したトラフィックが選定したターゲットリンクを通過するような、ターゲットエリアの周辺に存在する複数のデコイサーバを選択する。そして攻撃フェーズにおいては、多数のボットから複数のデコイサーバに対しトラフィックを流すことでターゲットリンクを高負荷とする [2]。

筆者らはこれまでに、CFA に先立ち攻撃者がターゲットリンクを選定するため大量の traceroute を行うことに着目し、traceroute の発生間隔に基づく攻撃ホストの検知法を提案した [8]。しかし [8] を含め既存の CFA の研究は、主に準備や攻撃段階における CFA を検出することを目的としており、CFA に対して脆弱なエリアを事前に予測し、重点的な設備投資などで事前に CFA に備える研究は見られない。

2. 節で述べるように、これまでに CFA を行うボットを識別し、探索フェーズもしくは攻撃フェーズにおいてパケットを規制することで CFA を防御する様々な方式が検討されている。しかし防御効果を向上させるためには CFA の発生を未然に防ぐ必要があり、攻撃フェーズ前で検知・防御を行うことが望ましい。ただし現時点ではネットワークのトポロジレベルで潜在的な CFA のターゲットを予測する方法は見られない。そこで本稿では、将来 CFA のターゲットとなる可能性の高いエリアに対してより効果的な防御を展開することを目的として、CFA に対して脆弱なエリアの高速で効率的な選択アルゴリズムを提案する。そして提案方式で選択されたエリアと理想的なエリア選択の結果を比較することで、提案アルゴリズムの性能を評価し、提案方式の有効性と精度を検証する。本アルゴリズムにより脆弱なエリアを特定することで、将来的にはこれらのエリアに対する防御をより効果的に展開することを目指す。以下、2. 節で関連研究について述べ、3. 節で CFA について述べ、4. 節で CFA の影響度の評価尺度について述べる。そして 5. 節で提案アルゴリズムの概要、6. 節で性能評価を行い、7. 節で全体をまとめる。

2. 関連研究

既存の CFA の検知や防御技術に関しこれまでに提案されている方式は、主に攻撃フェーズにおける技術と、探索フェーズにおける技術のいずれかに分類できる。攻撃フェーズにおける技術として、まず CFA の発生を検知するものがある [3] [4]。Narayanadoss らはターゲットリンクのトラフィック量を測定し、ANN, CNN, LSTM などの深層学習を用いて CFA の発生リンクを検知する方式を提案している [3]。Xue らはルータ間で E2E もしくはホップ間のアクティブ測定を行い、CFA の発生リンクを検知する方式を提案している [4]。ただしこれらの方式では CFA を行っている攻撃フローは特定できないため、CFA の防御は行えない。

さらに攻撃フェーズにおける技術として、Software-Defined Networking (SDN) を用いて迂回制御等で攻撃を防御する方式が提案されている [5] [6]。例えば Hyder らは ONOS Rest API と DNS のポートリダイレクションを利用したインテントベースのトラフィック変更を活用することで、CFA 防御の安全性を確保するフレームワークを提案している [5]。Rafique らはリンク選択、攻撃検知、悪意あるフローの遮断モジュールを採用した、CFADefense と呼ばれる新しい CFA 対策の設計と実装を提案している [6]。Aydeger らは CFA 攻撃を防御するための SDN ベースの MTD メカニズムを提案している。評価結果として、経路変異がネットワークサービスに大きな混乱を引き起

こすことなく、ターゲットリンクの負荷を効果的に削減できることを示している [7]。しかしこれらの方式では攻撃が発生してから検知・防御を行うので、CFA を未然に防ぐことはできない。

DDoS 攻撃に対する防御は、トポロジの観点から実行可能であると主張する研究も見られる [9] [10]。Guo らは、トポロジとトラフィック特徴に基づく深層学習手法 (GLD-Net) を提案している。トポロジとトラフィック特徴を初めて融合し、グラフニューラルネットワークを使用して高性能な DDoS 攻撃の侵入検出を実現する [10]。Liaskos らは、トポロジと検出効率の関係についての形式的な証明と、両者の関係を定量化する新しいオフライン測定手法を提案している。その結果、新たな評価基準がトポロジの検出関係を効果的に表現できることを示し、既存の広く用いられている評価基準は十分に本目標を果たせないことを明らかにしている [9]。しかし CFA を防御する際には、CFA の特徴を考慮する必要がある。

3. Crossfire Attack (CFA)

3.1 CFA の攻撃手法

CFA では探索フェーズと攻撃フェーズの 2 つのフェーズで実施される。探索フェーズにおいて、図 1 に示すように攻撃者は多数のボットホストからターゲットエリア (TA) 内の公開サーバや TA 周辺のデコイサーバに traceroute パケットを送信し、得られた経路情報をもとに攻撃リンクを選択する。ノードまでの経路情報を取得するツールである traceroute を実行することで、実行したノードから指定したノードまでの経路 (経由するルータ) のリストを得る。攻撃フェーズでは、攻撃者は多数のボットからデコイサーバに少量のトラフィックを生成し、TA と外部との通信を遮断する。

CFA における攻撃者の流れは以下になる。

- (1) 多数のボットから TA 内の多数の公開サーバ (Web サーバ等) に対し traceroute を行うことで TA 内のホスト宛てのフローの多くが経由する少数のリンク (攻撃対象リンク) を発見
- (2) 多数のボットから TA の周囲に存在する多数の公開サーバ (デコイサーバ) に対し traceroute を行い、攻撃対象リンクをフローが経由するボット・デコイサーバ組を選定
- (3) 選定したボット・デコイサーバ組に、検知されない程度の少量のトラフィックを生成 (通常の HTTP request/response など)

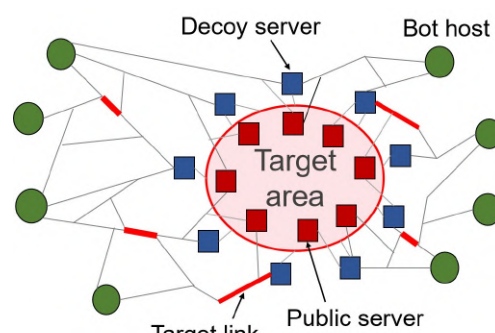


図 1 Crossfire Attack

3.2 CFA の検知の困難さ

CFA は、以下のような理由から検知・防御が困難である。

- 既存の CFA の防御策では、CFA の TA を予測することができないため、CFA 攻撃を迅速に検知し、効果的な防御を行うことができない。

- TA は直接攻撃を受けておらず、TA 近くのデコイサーバは異常なトラフィックを発見しないため、TA 内のサーバは CFA 攻撃を検知できない。

- CFA では大量の低速の攻撃フローがターゲットリンクを通過する。そのため、ターゲットリンクに接続されたルータは、攻撃フローと正当なフローの識別が難しい。

3.3 CFA の特徴

CFA では攻撃に先立ち、攻撃者は攻撃に用いるボット・デコイサーバ組を選定するため、大量のボットと TA 内のサーバ/デコイサーバ間に traceroute が発生する。この traceroute はネットワークのリンクが更新される前に選定を終える必要があるため、短い時間内に連続して行われる。また多くの場合、攻撃者はボットマーケット (PPI: pay-per install) を使用するが、コストは使用するボットの数に比例するため、1つのホストが多数のターゲット/デコイサーバに対し traceroute を実施する。この特徴をもとに、次節ではトラフィックのルーティング頻度に基づく CFA 攻撃に対するエリアの脆弱性を定義し、CFA の TA 選択アルゴリズムを提案する。

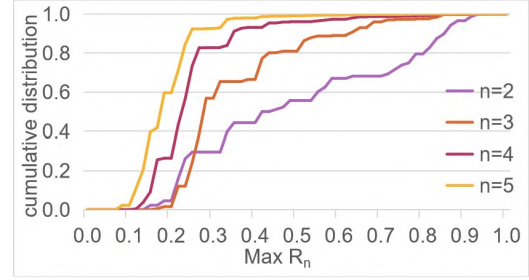
4. CFA の影響度の評価尺度

CFA の攻撃者は少量のトラフィックを大量の Bot から TA 周辺のデコイサーバ間で転送することで、目標エリア周辺に存在するリンクを高負荷にし、TA とそれ以外のエリア間のトラフィック流通を妨害する。そのため攻撃者は TA と外部との間のトラフィックの多くが通過するリンクを攻撃目標として選定する。ネットワークの隣接する複数のノードを CFA の TA x としたときに、 x と x 以外の領域とを跨ぐ少数のリンクを削除したときに、 x の外部との間のトラフィック量のうち通信不能となるものの割合が高いエリア x ほど、CFA に対して脆弱なエリアと考えられる。そのため CFA に対する脆弱性を測る尺度として、以下の変数を定義する。

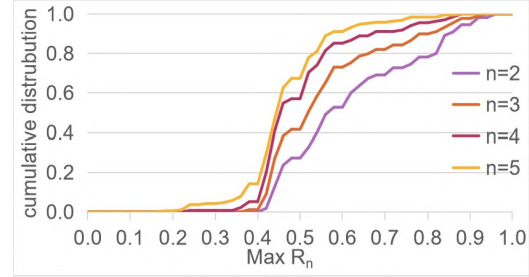
- (1) $A_n(x)$: n 個の隣接ノードで構成されるエリア x
- (2) $E_n(x, y)$: 任意の $A_n(x)$ に対して、 $A_n(x)$ と他エリアを跨る任意のリンク y
- (3) $R_n(x, y)$: $A_n(x)$ 以外の任意のノードと、 $A_n(x)$ の任意のノードとの間の最短ホップ経路のうち、リンク $E_n(x, y)$ を通るものの割合
- (4) $Max R_n(x)$: $R_n(x, y)$ の最大値
- (5) $Max_2 R_n(x)$: $R_n(x, y)$ の最大値と 2 番目に大きな値との合計値

図 2 に、図 5(a)(b) にトポロジを示す 2 つの米国の商用 ISP のネットワークにおいて各々、任意の n 個の隣接ノードからなるエリア x に対する $Max R_n(x)$ の累積分布を示す。 $Max R_n(x)$ が 90% を超えるエリアを図 3 に赤色で示すが、本エリアは複数のノードが直列に接続した構造であることがわかる。直列構造では、各ノードは前後の隣接ノードにのみ接続されており、そのエリアへのリンクは 2~3 本しかなく、トラフィックがそのうちの 1 本のリンクを通過して送信される場合が非常に多い。一方で n が大きくなると、ターゲットエリア内に含まれるノード数が増え、直列の部分が少なくなるため、 $Max R_n(x)$ は減少する。したがって、直列構造はネットワークのトポロジにおいて相対的に脆弱な部分であり、容易に CFA の対象となる。ごく少数のエリアのみが $Max R_n(x)$ の値が大きく、防御策を展開する際には、これら少数の CFA に脆弱なエリアに重点を置くことが重要である。

ただし CFA では $A_n(x)$ に隣接しないリンクを攻撃対象とすることも考えられるが、本稿では便宜上、 $A_n(x)$ と外部のエリアとの境界に存在するリンクを攻撃対象の候補として考える。



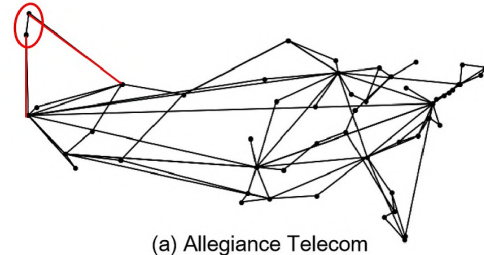
(a) Allegiance Telecom



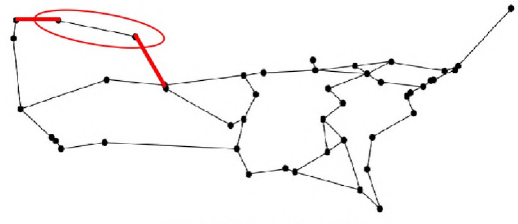
(b) At Home Network

図 2 Cumulative distribution of $Max R_n(x)$

各 $A_n(x)$ には対応する $Max R_n(x)$ と $Max_2 R_n(x)$ が存在し、CFA の攻撃者は $A_n(x)$ の CFA を行う際に、これらリンクを攻撃対象として選択することで、効率的・効果的に CFA を行うことが可能となる。



(a) Allegiance Telecom



(b) At Home Network

図 3 Examples of serial structure in topologies

5. TA 選択アルゴリズム

本稿ではトポロジの特徴に基づいて CFA に対して脆弱な TA を抽出するアルゴリズムを提案する。多くの場合、CFA の TA は多くのノードを含まないため、本稿では TA を z 個の隣接ノードから構成されるエリアとする。理想的には、トポロジ上の隣接する z 個の全てのエリア x について $Max_2 R_n(x)$ を計算し、その値が与えられた閾値 T より大きな全てのエリアを抽出できればよい (理想法)。しかし $Max_2 R_n(x)$ の計算には全てのノード間の最短ホップ経路が必要であるため、ノード数が多い大規模ネットワークでは計算量が大きくなる。そこで、より少ない計算量で効果的に $Max_2 R_n(x)$ が大きなエリアを抽

出するアルゴリズムを検討する。

ネットワークトポロジ内の全ノードの中で、上位 30% 以内の次数を有するノードを高次数ノードとして定義する。TA 内の高次数ノード数の増加に伴い、TA 内のノードの平均次数は増加し、TA と外部との接続リンク数が増加する。そのため TA 外から TA 内に送信されるトラフィックは多数のリンクに分散しやすくなり、 $Max R_n(x)$ や $Max_2 R_n(x)$ は減少する。したがって、TA 内の高次数ノードの数が少ないほど TA は CFA に対して脆弱になる。

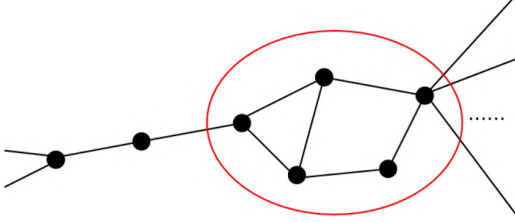


図 4 Example of quasi-serial structure in topology of Allegiance Telecom

また次数が 2 のノードが直列に連結した直列構造となったエリアの $Max_2 R_n(x)$ は大きなことが予想される。さらに提案選択アルゴリズムでは、直列構造に加えて准直列構造の判定を追加する。ただし准直列構造を次数が 3 以上のノードが複数連結したもので、エリア内外が 2 個のノードでのみ連結しているエリアと定義する。Allegiance Telecom の一部のエリアを図 4 に示すが、赤色で示すエリアが准直列構造の例である。

これらの特徴から、各エリア x に対して以下で定義するスコア $S(x)$ を計算し、その値が大きなエリアを CFA に対して脆弱なエリアとして抽出する。

$$S(x) = s_1(x) + s_2(x) + s_3(x) \quad (1)$$

ただし $s_1(x)$ をエリア x と隣接する直列構造の数、 $s_2(x)$ をエリア x と隣接する准直列構造の数、 $s_3(x)$ をエリア x 内の高次数ノード数の逆数と各々定義する。

提案アルゴリズムは各ノードの次数のみで計算できるため、最適法と比較して計算時間を大幅に低減することが可能である。

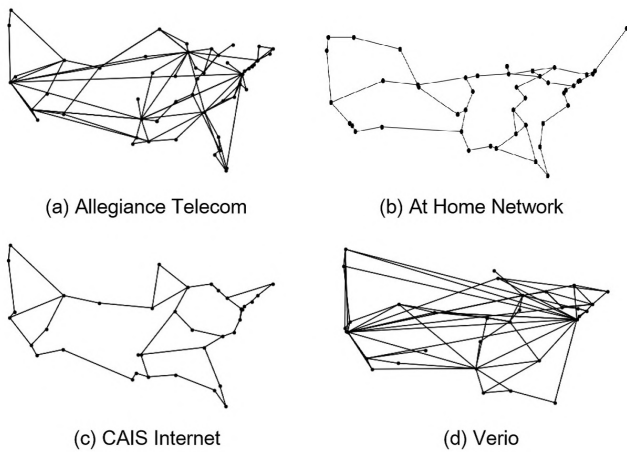


図 5 Topologies of four networks

6. 性能評価

6.1 評価条件

米国の商用バックボーン ISP である At Home Network, Al-

legiance Telecom, CAIS Internet, Verio の 4 つのネットワークのトポロジを評価に用いる。図 5 に、これら 4 つのネットワークのトポロジを図示する。また考慮するエリアのノード数 z を 5 とする。

6.2 有効性評価

提案アルゴリズムで発見した各エリアの $Max_2 R_n(x)$ を最適法の発見エリアと比較することで提案アルゴリズムの有効性を確認する。ただし最適法では、 $Max_2 R_n(x)$ が、任意に与えた閾値 T 以上となるエリアを全て選択した。一方、提案アルゴリズムではスコア $S(x)$ が大きな順に、最適法で発見された個数と同じ個数のエリアを選択した。

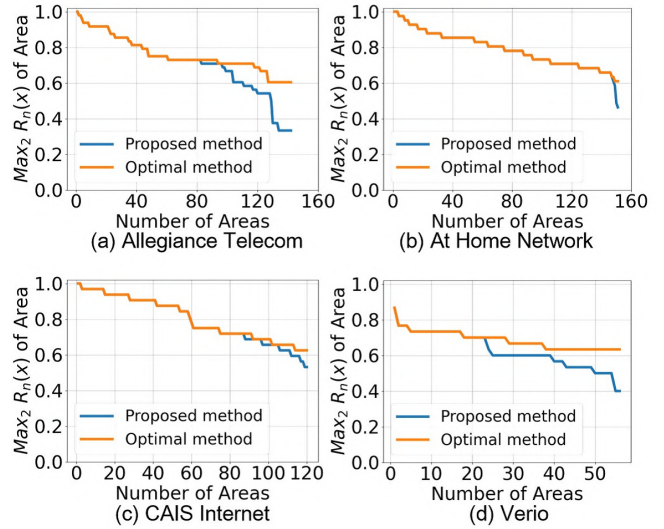


図 6 $Max_2 R_n(x)$ of each area selected by each method in descending order when setting $T = 0.6$

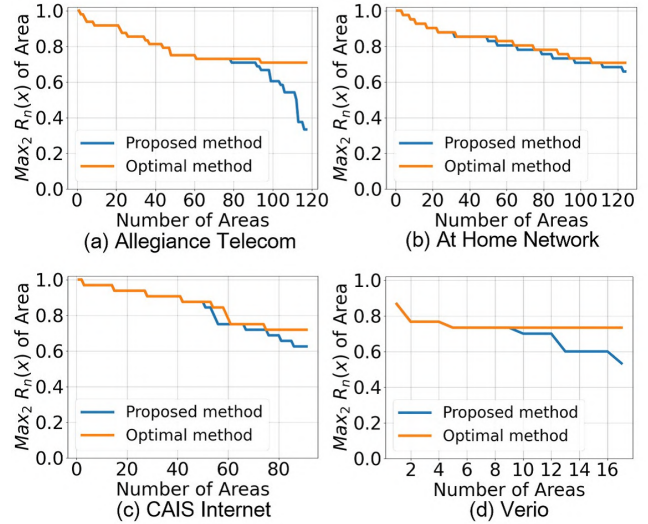


図 7 $Max_2 R_n(x)$ of each area selected by each method in descending order when setting $T = 0.7$

図 6 に、最適法の閾値を $T = 0.6$ に設定したときの、各方式で選択されたエリアの $Max_2 R_n(x)$ を降順にプロットする。同様に図 7、図 8、図 9 に最適法の閾値を $T = 0.7$, $T = 0.8$, $T = 0.9$ に設定した場合の結果を各々示す。ただし Verio においては、0.8 以上の $Max_2 R_n(x)$ を有するエリアが存在しなかったため、図 8 と図 9 には Verio の結果は除いている。

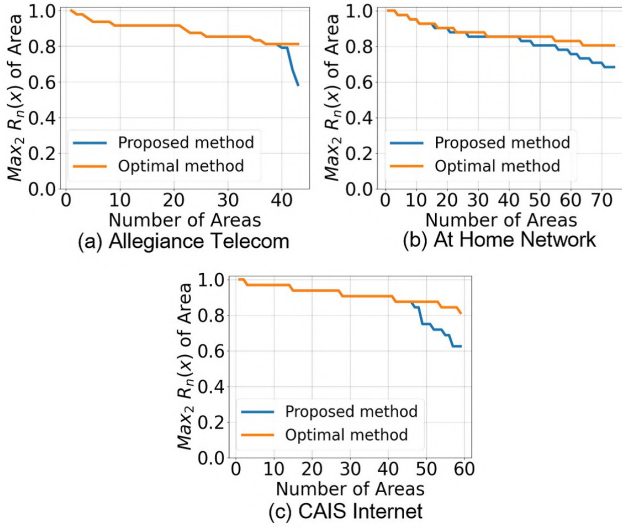


図 8 $Max_2 R_n(x)$ of each area selected by each method in descending order when setting $T = 0.8$

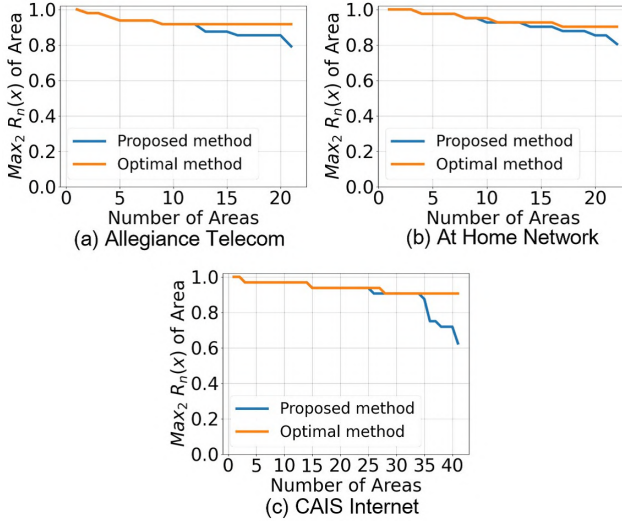


図 9 $Max_2 R_n(x)$ of each area selected by each method in descending order when setting $T = 0.9$

図 6(a) に示すように Allegiance Telecom で $T = 0.6$ の場合、上位 80 個程度の $Max_2 R_n(x)$ の値が大きなエリアは提案方式で正しく抽出できるが、残る 60 個程度のエリアについては抽出精度が低下する傾向がある。これは Allegiance Telecom のような複雑なネットワークポロジで、特に $Max_2 R_n(x)$ の値が 0.8 程度未満の場合、ノード次数のみでは脆弱なエリアとそれ以外のエリアを正しく区別することが難しいためである。しかし提案アルゴリズムは $Max_2 R_n(x)$ が 0.8 以上といった特に CFA に対して脆弱なエリアを正確に抽出する。

6.3 提案方式の抽出精度

図 6～図 9 より、提案アルゴリズムは評価に用いた 4 つのネットワークポロジにおいて、高い精度で CFA に対して脆弱なエリアを抽出でき、特に $Max_2 R_n(x)$ が 0.9 以上の場合、高い特性精度が得られることを確認した。しかし $Max_2 R_n(x)$ が 0.6 以下のエリアに対しては誤検出が大きい。本節では提案方式の抽出精度をより詳細に分析する。

提案方式の抽出精度を測る尺度として、ここでは、 $Max_2 R_n(x)$ が T 未満 (抽出すべきでないエリア) の中で、誤って抽出されたエリアの割合 (FPR: false positive ratio) と、 $Max_2 R_n(x)$

が T 以上 (抽出すべきエリア) の中で、誤って抽出されなかったエリアの割合 (FNR: false negative ratio) を用いる。

図 10～13 に、4 つの各ネットワークポロジにおいて、閾値 T の 4 つの各値における、提案アルゴリズムの FPR と FNR を示す。やはり図 13 に示す Verio の結果では、 $T = 0.8$ と $T = 0.9$ の場合は該当するエリアが存在しないため、これらの結果は省いている。

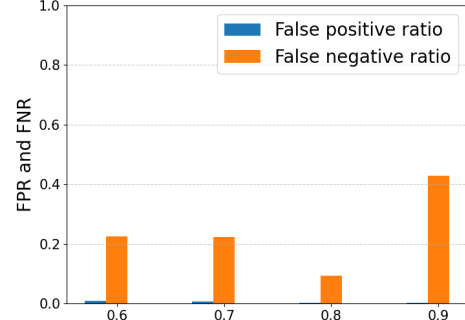


図 10 FPR and FNR in Allegiance Telecom

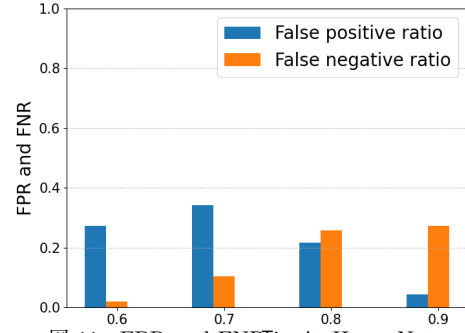


図 11 FPR and FNR in At Home Network

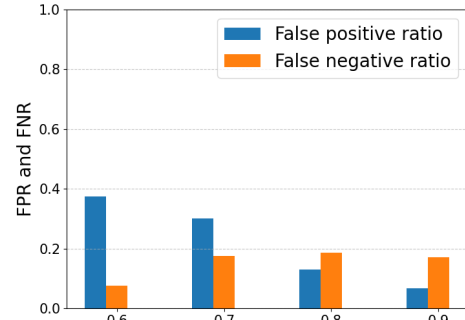


図 12 FPR and FNR in CAIS Internet

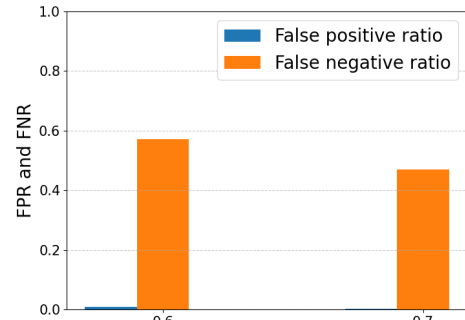


図 13 FPR and FNR in Verio

Allegiance Telecom と Verio は他の 2 つのネットワークと比較して FNR が高い傾向が見られる。 $T = 0.9$ の Allegiance Telecom と、 $T = 0.6$ と 0.7 の Verio の FNR は 0.4 以上である。一方で、Allegiance Telecom と Verio の FPR は常に低くゼロに近い値であり、提案アルゴリズムは $Max_2 R_n(x)$ が設定した閾値 T より大きいエリアを漏れなく抽出できる。しかし Allegiance Telecom や Verio などの複雑なネットワークトポロジでは、提案アルゴリズムは多くの CFA に脆弱なエリアを見逃すことがわかる。

At Home Network と CAIS Internet では逆の結果が示されている。両者とも FNR は比較的低いが、FPR は $T = 0.6, 0.7$ のとき 0.2 以上である。これら 2 つのネットワークトポロジにおいては、 $Max_2 R_n(x)$ の値が 0.6 より小さいエリアが多く存在するため、これら単純なネットワークトポロジでは、提案アルゴリズムは条件を満たさないエリアを誤って選択する可能性が高い。

4 つのトポロジの FPR が全体的に閾値 T の増加とともに減少する傾向が見られる。各ネットワークトポロジにおいて、 $Max_2 R_n(x)$ の値が 0.4~0.6 の範囲で多くのエリアが存在するため、小さな閾値 T (例えば 0.6 や 0.7) の場合、提案アルゴリズムはこれらのエリアを誤抽出する可能性が高い。一方、閾値 T が増加すると、これら T の小さなエリアの誤抽出が減少し、提案アルゴリズムの FPR が減少する。At Home Network と CAIS Internet で特にこの傾向が顕著であり、これらの 2 つのネットワークトポロジが比較的単純であり、0.4~0.6 の範囲での $Max_2 R_n(x)$ の値が多いことが原因と思われる。

7. ま と め

CFA はネットワークのリンクを攻撃対象とする性質のため、従来の DDoS の防御技術は適用が困難である。CFA に対して多くの防御法が検討されているが、CFA が選択するターゲットリンクやエリアをどのように予測するかという問題は未解決である。CFA を未然に効果的に防御するためには、CFA に対して脆弱なネットワーク上のエリアを抽出し、そのようなエリアに対して重点的に設備増設を行うことが有効と思われる。そこで本稿では、ネットワークトポロジの構造から、CFA に対して脆弱なエリアを測るための指標を定義した。そして本尺度が大きい、CFA に対して脆弱なエリアを少ない計算量で高精度に抽出する脆弱エリア抽出アルゴリズムを提案した。そして 4 つの米国商用 ISP のネットワークトポロジを用いた数値評価により、提案アルゴリズムがネットワークトポロジの CFA に対し脆弱なエリアを高精度に抽出できることを明らかにした。

提案アルゴリズムはネットワークトポロジの CFA に対し脆弱なエリアを高精度に効率的に識別できるものの、抽出精度にはまだ改善の余地がある。そのため今後はネットワークトポロジの複雑さを表す新たな指標を導入し、複雑さの異なるネットワークトポロジに対して異なる閾値 T を設定し、対応するトポロジに適応させることで、異なるネットワークトポロジごとにアルゴリズムの抽出精度を向上させる予定である。

謝辞 本研究成果は JSPS 科研費 21H03436 と 21H03437 の助成を受けたものである。ここに記して謝意を表す。

文 献

- [1] C. Dietzel, M. Wichtlhuber, G. Smaragdakis, and A. Feldmann, Stellar: Network Attack Mitigation using Advanced Blackholing, ACM CoNEXT 2018.
- [2] M. Kang, S. B. Lee, and V. D. Gligor, The Crossfire Attack, IEEE SSP 2013.
- [3] A. R. Narayanadoss, M. Gurusamy, P. M. Mohan, and T. Truong-Huu, Crossfire Attack Detection using Deep Learning in Software Defined ITS Networks, VTC Spring 2019.
- [4] L. Xue, E. W. W. Chan, G. Gu, X. Luo, X. Ma, and T. T. N. Miu, LinkScope: Toward Detecting Target Link Flooding Attacks, IEEE Trans. Information Forensics and Security, 2018.
- [5] M. F. Hyder and T. Fatima, Towards Crossfire Distributed Denial of Service Attack Protection Using Intent-Based Moving Target Defense Over Software-Defined Networking, IEEE Access, vol. 9, pp. 112792-112804, 2021.
- [6] W. Rafique, X. He, Z. Liu, Y. Sun, and W. Dou, CFADefense: A Security Solution to Detect and Mitigate Crossfire Attacks in Software-Defined IoT-Edge Infrastructure, IEEE HPCC/SmartCity/DSS 2019.
- [7] A. Aydeger, N. Saputro, K. Akkaya, and M. Rahman, Mitigating Crossfire Attacks Using SDN-Based Moving Target Defense, IEEE LCN 2016.
- [8] Manami Nakahara and Noriaki Kamiyama, Detecting Crossfire-Attack Hosts in Search Phase, APNOMS 2022 (Poster Session).
- [9] C. Liaskos and S. Ioannidis, Network Topology Effects on the Detectability of Crossfire Attacks, IEEE Transactions on Information Forensics and Security, vol. 13, no. 7, pp. 1682-1695, July 2018.
- [10] Wei Guo, Han Qiu, Zimian Liu, Junhu Zhu, and Qingxian Wang, GLD-Net: Deep Learning to Detect DDoS Attack via Topological and Traffic Feature Fusion, Computational Intelligence and Neuroscience, vol. 2022, Article ID 4611331, 2022.