

GCNを用いた独自Fake型CPAに対する脆弱ノード検出法

上野 優衣[†] 上山 憲昭[†]

[†] 立命館大学 情報理工学部

〒567-8570 大阪府茨木市岩倉町 2-150

E-mail: †is0547ph@ed.ritsumei.ac.jp, ††kamiaki@fc.ritsumei.ac.jp

あらまし CDN (content delivery network) に代わるネットワークとして, ICN (information-centric networking) の研究が進められている. しかし, ICN には悪意を持ったユーザが不当なコンテンツをルータにキャッシュさせることで, キャッシュヒット率を低下させるコンテンツポイズニング攻撃 (CPA: content poisoning attack) という脆弱性が存在することが指摘されている. 中でも, 独自 Fake 型 CPA は検知がしづらく, その対策についての研究は十分になされていない. そこで, 本稿では畳み込みグラフニューラルネットワーク (GCN: graph convolutional network) を用いて, Fake コンテンツが注入される可能性の高い脆弱なノードを検出する手法を提案する. さまざまな条件下での提案方式の予測精度を評価し, 有効な推測を行うことが可能な教師データ数の閾値を明らかにする.

キーワード ICN, コンテンツポイズニング攻撃, GCN

Vulnerable Node Detection Method for Unique Fake-type CPA using GCN

Yui UENO[†] and Noriaki KAMIYAMA[†]

[†] College of Information Science and Engineering, Ritsumeikan University

2-150 Iwakura-cho, Ibaraki, Osaka 567-8570

E-mail: †is0547ph@ed.ritsumei.ac.jp, ††kamiaki@fc.ritsumei.ac.jp

Abstract Information-centric networking (ICN) is being studied as an alternative to content delivery networks (CDNs). However, it has been pointed out that ICN is vulnerable to a content poisoning attack (CPA), in which a malicious user caches inappropriate content in the router, thereby decreasing the cache hit ratio. In particular, original Fake CPAs are difficult to detect, and there has been insufficient research on countermeasures against them. In this paper, we propose a method to detect vulnerable nodes that may be injected with fake content using a convolutional graph neural network (GCN). We evaluate the prediction accuracy of the proposed method under various conditions and investigate a threshold for the number of supervised data at which effective guesses can be made.

Key words ICN, Content Poisoning Attack, GCN.

1. はじめに

現在, コンテンツ配信のプラットフォームとして Content Delivery Network (CDN) が主に使用されている. CDN では, コンテンツをサーバにキャッシュすることで, 高速なコンテンツ配信を実現している. 近年, スマートフォンなどの高性能な移動型端末の普及により, コンテンツプロバイダによる映画やドラマなどの大容量コンテンツの配信サービスの需要がさらに拡大している. また, スマートフォンの普及に伴う SNS や動画投稿サービスなどのユーザが生成するコンテンツの投稿サービスの普及も, インターネットのトラフィックが増加する原因となっている. インターネットでは, IP アドレスをもとにパケット転送を行うホスト指向の通信が行われているが, コンテンツ配信がトラフィックの大部分を占

める昨今では, ホスト指向の通信方式は非効率的である.

そこで, コンテンツ名を指定してコンテンツを要求することで, 従来の IP で必要であった名前解決のためのオーバヘッドを削減可能かつ, 要求パケット (Interest) の転送経路上のルータからもコンテンツの取得を可能とする情報指向ネットワーク (ICN: information-centric networking) の研究が進められている [1]. また, ICN では誰でもコンテンツをアップロードすることができるため, 近年のユーザ生成コンテンツの増加という傾向にも適応している.

しかし, 悪意を持ったユーザが不当なコンテンツをルータにキャッシュさせることで, キャッシュヒット率を低下させるコンテンツポイズニング攻撃 (CPA: content poisoning attack) の可能性が指摘されている [2]. CPA によって, ルータのキャッシュ (CS: content store) に偽りのコンテンツが

キャッシュされキャッシュヒット率の低下などを引き起こす問題がある。

ICN には、コンテンツの要求者が受信コンテンツに対して、コンテンツに対応する公開鍵で生成されたデジタル署名が一致するかを確認することでコンテンツの正当性を確認する方法が存在する。しかし、自身で無意味なオリジナルのコンテンツ (Fake コンテンツ) を生成し、Bot からコンテンツを要求させる独自 Fake 型 CPA [3] は、正常なユーザがコンテンツを受信することがないため、デジタル署名による検知が困難である。独自 Fake 型 CPA において、不当なコンテンツを排除するにはキャッシュされているコンテンツ全てを正当なコンテンツであるか検証する必要がある。しかし、ネットワーク上の全てのノードのキャッシュコンテンツを全て検証することは現実的ではない。そこで本稿では、Fake コンテンツが注入される可能性の高い脆弱なノードを、一部のノードのキャッシュ情報をもとに GCN (graph convolutional network) を用いて推測する方式を提案する。

2. 関連研究

CPA には大きく分けて 4 つの型が存在する [3]。

- **結託 Corrupted 型:** 実在するコンテンツを騙った署名が一致しない Fake コンテンツを Bot の要求に応じて配信
- **自然 Corrupted 型:** 実在するコンテンツを騙った署名が一致しない Fake コンテンツを正常ユーザからの要求に応じて配信
- **詐称 Fake 型:** 実在するコンテンツのデジタル署名を持った Fake コンテンツを正常なユーザからの要求に応じて配信
- **独自 Fake 型:** 攻撃者が独自に生成した Fake コンテンツを Bot の要求に応じて配信

Corrupted 型 CPA の 2 種類は署名が一致しない Fake コンテンツを配信するため、デジタル署名を検証することで容易に検知が可能である。また、検知した Fake コンテンツを CS から削除するために、検知した不当なコンテンツをルータに通知する [4] などの防御法が研究されている。

詐称 Fake 型 CPA は、公開鍵を管理する認証局 (CA: certification authority) の職員と攻撃者が結託して、実在するコンテンツの署名に使用している公開鍵を攻撃者の公開鍵と書き換えることによって成立する攻撃である。これは公開鍵の一つの機関でのみ管理することに脆弱性がある。これに対して、分散型台帳である IOTA でコンテンツ名を管理することで、詐称 Fake 型 CPA の発生を防御する手法が提案されている [5]。また、詐称 Fake 型 CPA は CA 職員との結託が必要であるため、頻繁に実現させることは難しいと考えられる。

一方で独自 Fake 型 CPA では、攻撃者が自身で生成した Fake コンテンツをルータに注入する。そのため Fake コンテンツを要求し、Interest を送付するのは Bot のみであり、正常ユーザからの要求は発生しないため検知されづらい。また、攻撃者が独自に生成したコンテンツそのものに署名を取得しているため、詐称 Fake 型のような CA 職員との結託も不要で実現が容易である。しかし、独自 Fake 型 CPA がネットワークに与える影響については調査が行われているが [3]、独自 Fake コンテンツを検知する手法は提案されていない。

また独自 Fake 型 CPA の Fake コンテンツは、Bot のみに依存する Interest 発生の特性から、Fake コンテンツのオリジナルが存在するノードまたは Bot が配置されているノードの近くにキャッシュされやすいことが明らかにされている [3]。そのため、ネットワーク内で Fake コンテンツがキャッシュされているノードが判明すれば、周辺ノードに Fake コンテ

ンツのオリジナルか Bot が存在する可能性が高いと考えられる。このように、ネットワークを構成するノードのトポロジ上の位置によって、独自 Fake 型 CPA に対する脆弱度合い、すなわち Fake コンテンツが注入される可能性は異なる。そのため事前に、Fake コンテンツが注入される可能性の高いノードが検知できれば、それら検知されたノードに対してのみ、重点的にキャッシュ容量を増設したり、キャッシュコンテンツの正当性を調査するなどの、効率的・効果的な対策をとることが可能となる。

本稿では、独自 Fake 型 CPA に対しての効果的な事前対策を可能とすることを目的に、Fake コンテンツが注入される可能性が高い脆弱なノードを検出することで、独自 Fake 型 CPA の防御が特に必要なノードを明らかにする方式を提案する。

3. 提案方式

3.1 GCN

従来のニューラルネットワークが画像やテキストを学習するモデルであるのに対し、グラフニューラルネットワーク (GNN: graph neural network) はグラフ構造データを学習させる機械学習モデルである。GNN はさまざまな分野において、その活用が研究されている [6]。[7] では、ネットワーク侵入検知システム (NIDS: network intrusion detection systems) における GNN の有効性が検証されており、ネットワークセキュリティ分野においても GNN の適用が可能と考えられる。

グラフ畳み込みニューラルネットワーク (GCN) は、GNN の中でも代表的な手法のひとつであり、各ノードと隣接するノードの特徴量を畳み込むことで学習を行う。GCN には、ノードのクラス分類、ノード間にリンクが存在するかどうかの予測、グラフのクラス分類などさまざまなタスクが存在する [8]。

ノードのクラス分類のタスクは、グラフ、全ノードのそれぞれの特徴量、一部のノードが属するクラス (ラベル) が入力として与えられた際に、半教師あり学習を行うことで、ラベルが未判明である残りのノードのラベルを予測するタスクである。GCN では、グラフの接続されたノード同士は同一ラベルを持つ可能性が高いという仮定のもと、伝播規則 [9] に基づいて特徴量を学習層ごとに更新し、最終的な予測ラベルを確定する。

3.2 GCN で使用する条件

Fake コンテンツが注入される可能性のある脆弱なノードを推測するため、GCN のラベルは、各ノードの CS における Fake コンテンツの平均注入率 (AIRF: average injection ratio of fake content) を元に作成する。AIRF を算出するには、各ノードの CS の全てのコンテンツに対して、コンテンツが正常コンテンツであるか Fake コンテンツであるかを確認し、キャッシュされている Fake コンテンツの数を調べる必要がある。そのため、少ない正解ラベルから脆弱なノードを効率よく推測できることが望ましい。また、全ノードの特徴量として、ネットワーク内で人気の高い上位 5 つのコンテンツについて、各ノードにおけるキャッシュヒット率を利用する。

4. 評価条件

4.1 教師データ作成条件

計算機シミュレーションによって、GCN で使用する特徴量とラベルを作成する。シミュレータには、米国の商用 ISP のバックボーンネットワークトポロジである At Home Network

と Allegiance Telecom を使用する。キャッシュ方式は、転送経路上の全てのルータでキャッシュする All Cache 方式とし、キャッシュ容量は 100(コンテンツ) とする。

正当コンテンツ数は $M = 10,000$ とし、Fake コンテンツ数 F は 32 または 1,024 とする。正常ユーザは $\theta = 0.8$ の Zipf 分布に従い正当コンテンツを要求し、Bot の要求比率 ρ を、単位時間あたりに生成される全正常要求数に対する、単位時間あたりに Bot から生成される CPA の要求数と定義し、数値評価では $\rho = 0.1, 0.3, 0.5$ の三つの場合を評価する。シミュレーション時間は 10,000 秒とし、1,000 秒に 1 回、各ノードの Fake コンテンツの注入率を測定し、測定した 10 回の平均値を各ノードの AIRF とする。Fake コンテンツのオリジンサーバと Bot は、それぞれ同数 N_{attack} 個を配置する。数値評価では、 N_{attack} が 1, 2, 4 の三つの場合を評価する。Fake コンテンツオリジンサーバと Bot の配置方式は、次の 3 種類とする。

- **Random 方式:** Fake コンテンツのオリジンサーバと Bot をランダムに配置
- **DCf 方式:** 次数中心性 (DC: Degree Centrality) が最大の N_{attack} のノードに Fake オリジンサーバを先行配置し、続いて、これら Fake オリジンサーバの配置ノードから、最も遠いノードに Bot を後続配置 [3]
- **DCb 方式:** DC に基づいて、Bot を N_{attack} 箇所に先行配置して、続いて、これら Bot 配置ノードから、最も遠いノードに Fake オリジンサーバを後続配置 [3]

AIRF を元に、10,000 秒の間に一度でも Fake コンテンツが CS に注入された形跡のあるノード ($AIRF > 0$) のクラスを Label1 とし、一度も Fake コンテンツが注入されていないノード ($AIRF = 0$) のクラスを Label0 と定義する。

4.2 GCN の学習条件

グラフには、教師データの作成に使用したのと同じ Allegiance Telecom と At Home Network のネットワークポロジの接続情報を与える。全ノードの特徴量として、人気の高い上位 5 つのコンテンツの各ノードにおけるキャッシュヒット率を使用する。ただしキャッシュヒット率は計算機シミュレーションにより算出する。

また GCN の作成には、Python のライブラリである PyTorch を使用する。PyTorch ではモデルの重みが実行のたびにランダムに初期化されるため、実行のたびに結果が変化し、予測精度に影響を及ぼす。そこで seed 値を設定することで、モデルの初期化にばらつきが出ないように固定する。2 種類の seed 値 (42, 1,024) を設定して、精度の良いものを研究結果として用いる。

5. 性能評価

5.1 評価尺度

偽陰性率 (FNR: false negative ratio), 偽陽性率 (FPR: false positive ratio), 精度の 3 つの尺度を用いて性能を評価する。表 1 に示すように、TP(TN) を Label1(0) のノードを正しく判定した数とし、FN(FP) を Label1(0) のノードを誤って判定した数とする。

表 1: 陽性陰性の定義

	正解 Label1	正解 Label0
予測 Label1	TP	FP
予測 Label0	FN	TN

FNR は Label1 のノードを誤って Label0 と判定する割合で、FPR は Label0 のノードを誤って Label1 と判定する割合である。FPR, FNR を以下のように定義する。

$$FPR = \frac{FP}{FP + TN} \quad (1)$$

$$FNR = \frac{FN}{TP + FN} \quad (2)$$

ネットワークポロジと Fake コンテンツや Bot の配置方式の関係により、Fake コンテンツが広範囲のノードの CS に注入される場合と少数ノードの CS にのみに注入される場合がある。少数のノードの CS に Fake コンテンツが注入されている場合、精度が極端になることが考えられる。そのため、本稿では正解データの Label1 の割合がネットワークポロジ全体のうち 15% 以上となる配置が行われたケースについて比較し、考察を行う。

また、ラベルの教師データを以下の 2 方式で選択したときの予測精度を比較し、教師データの与え方が予測精度に与える影響を分析する。

- **無差別選択方式:** 教師データ数 X_{train} が偶数の場合は各ラベルが均等な数になるようランダムに選択し、奇数の場合は Label0 が Label1 の教師数より 1 つ多くなるように選択したラベルの教師データを与える。
- **次数交錯方式:** ネットワークポロジの次数中心性が最も高いものと最も低いものから交互に選択

無差別選択方式では、ラベル選択におけるランダム性の観点から、予測精度が極端に高い、または低くなる可能性が考えられる。そのためラベル選択におけるランダムシードを 3 種類与え、3 種類の教師選択パターンの FPR, FNR の平均をとって計測を行う。また、2 方式間での FPR(FNR) の値を比較するため、無差別選択方式の FPR(FNR) から次数交錯方式の FPR(FNR) を減算したものを差異と定義する。FPR(FNR) の値が小さいほど正しく予測が行えていることを表すため、差異が正の場合は次数交錯方式が優位な方式、差異が負の場合は無差別選択方式が優位な方式と判断する。

5.2 要求比率による差異

At Home Network において、Fake オリジンサーバと Bot の配置数 N_{attack} を 4 とし DCb 方式で配置し、Fake コンテンツ数が $F = 32$ であるときに、Bot の要求比率 ρ を変化させたときの予測精度を比較する。

ρ が 0.1, 0.3, 0.5 の全ての場合において、Label1 となるノードは 17 箇所である。教師データ数を 3 から 6 に変化させたときの、FNR, FPR, それぞれの差異を表 2, 3, 4 に示す。全ての要求比率において、FPR の差異は負、FNR の差異は正である。このことから、無差別選択方式は FPR を抑えるのに有効で、次数交錯方式は FNR を抑えるのに有効であることが推測される。一方で、差異は教師数が増えると低くなっているが、これはどちらの方式においても教師数が増えると FNR が減少するため、差が小さくなっていると考えられる。本稿では、Fake コンテンツの注入されやすい脆弱なノードを推測することを目的としているため、FNR をより抑えられることが望ましく、次数交錯方式が優位といえる。

5.3 Fake コンテンツ数による差異

At Home Network において、Fake オリジンサーバと Bot の配置数 N_{attack} を 1 とし Random 方式で配置し、Bot の要求比率が $\rho = 0.3$ であるときに、Fake コンテンツ数 F を変化させたときの予測精度を比較する。

F が 32, 1,024 の場合において、Label1 となるノードはどちらも 10 箇所である。教師データ数を 3 から 6 に変化させ

表 2: $\rho = 0.1$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0370	0.9375	0.2667	0.5357	-0.2296	0.4018
4	0.2716	0.8000	0.9333	0.0741	-0.6616	0.7259
5	0.2051	0.8444	0.5714	0.3333	-0.3662	0.5111
6	0.5769	0.5000	0.8571	0.2308	-0.2801	0.2692

表 3: $\rho = 0.3$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0370	0.8958	0.4667	0.3929	-0.4296	0.5029
4	0.2592	0.7777	0.9333	0.0741	-0.6740	0.7036
5	0.1666	0.7999	0.7143	0.1481	-0.5476	0.6518
6	0.5384	0.4762	0.9286	0.1154	-0.3901	0.3608

表 4: $\rho = 0.5$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0493	0.8958	0.4667	0.4643	-0.4173	0.4315
4	0.2345	0.8222	0.8667	0.0741	-0.6321	0.7481
5	0.1666	0.8222	0.7143	0.2222	-0.5476	0.6000
6	0.5641	0.4762	0.9286	0.1154	-0.3645	0.3608

たときの, FNR, FPR, それぞれの差異を表 5, 6 に示す. 無差別選択方式では, FPR と FNR のどちらかが低いともう一方が高くなる場合が多いが, 次数交錯方式では, FPR と FNR のどちらも低く抑えられている. そのため, FPR, FNR とともに次数交錯方式の効果が高くなる場合もあることが確認された.

FPR, FNR でどちらも差異が 10% 以上であった, $F = 32$ の教師数 4 の場合の, 各方式の予測トポロジと正解ラベルのトポロジを図 1 に示す. 多くのノードで正しくラベルが推定できている.

表 5: $F = 32$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.2843	0.4444	0.2500	0.4000	0.0343	0.0444
4	0.4902	0.3333	0.3750	0.1176	0.1152	0.2157
5	0.2020	0.7083	0.4286	0.1765	-0.2265	0.5318
6	0.4040	0.3333	0.4286	0.0606	-0.0245	0.2727

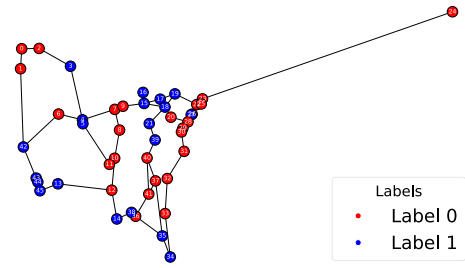
表 6: $F = 1,024$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0980	0.8889	0.2500	0.7143	-0.1519	0.1746
4	0.4019	0.4583	0.3750	0.2353	0.0269	0.2230
5	0.1414	0.9166	0.2857	0.3824	-0.1443	0.5342
6	0.3333	0.7142	0.4286	0.1515	-0.0952	0.5627

5.4 攻撃者配置数による差異

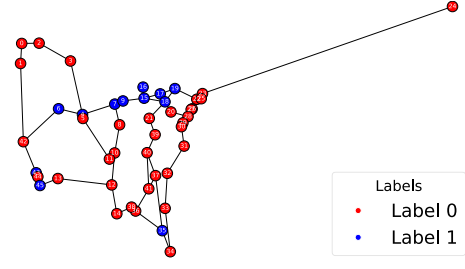
At Home Network において, DCb 方式で Fake コンテンツ数が $F = 32$ で, Bot の要求比率 $\rho = 0.3$ であるときに, Fake コンテンツオリジンサーバと Bot の配置箇所 N_{attack} を 1, 2, 4 箇所に変化させたときの予測精度を比較する.

攻撃者の配置箇所が 1, 2, 4 のとき, Label 1 となるノードはそれぞれ 8, 15, 17 箇所である. 教師データ数を 3 から 6 に変化させたときの, FNR, FPR, それぞれの差異を表 7, 8, 9 に示す. $N_{attack} = 2$ において, 次数交錯方式の FNR は低く抑えられているが, FPR がいずれの場合も 1 となって



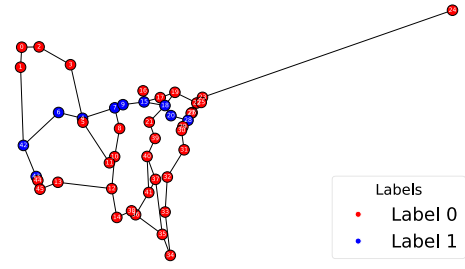
Predict : 無差別選択方式

(a) 無差別選択方式の予測ラベル



Predict : 次数交錯方式

(b) 次数選択方式の予測ラベル



Answer

(c) 正解ラベル

図 1: At Home Network, $N_{attack} = 1$, Random 方式配置, 要求比率 $\rho = 0.3$, 教師数 4 の GCN による予測と正解ラベル

いる. これは, 本来脆弱ではないノードを全て脆弱性があると判断しているため, FNR がいかに低くても有効とはいえない. また, 無差別選択方式において, $N_{attack} = 4$ のとき, $N_{attack} = 1, 2$ の場合と比べて, FPR の増加が見られる. 図 2 に, 各配置数での AIRF が大きい順に並べたものを示す. 攻撃者が 1, 2 のときに比べ, 攻撃者が 4 の場合は AIRF が 0.08 から 0.12 の間で横ばいになっている数が多い. 攻撃者を多数配置することで多くのノードに満遍なく Fake コンテンツがキャッシュされるため, Fake コンテンツをキャッシュさせることで各ノードに与える影響が小さくなる. その結果, キャッシュヒット率に与える影響が小さくなり, GCN を用いた脆弱ノードの推定精度が低下した可能性が考えられる.

表 7: $N_{attack} = 1$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.3703	0.3809	0.0000	0.7027	0.3703	-0.3217
4	0.5833	0.0555	0.1667	0.1111	0.4166	-0.0555
5	0.3619	0.3333	0.0000	0.6944	0.3619	-0.3610
6	0.5238	0.1333	0.2000	0.2286	0.3238	-0.0952

5.5 攻撃者の配置方式による差異

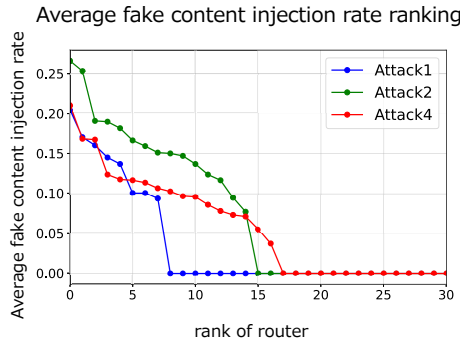
At Home Network において, Fake オリジンサーバと Bot

表 8: $N_{attack} = 2$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0459	0.9762	1.0000	0.0690	-0.9540	0.9072
4	0.3333	0.6666	1.0000	0.0000	-0.6666	0.6666
5	0.2500	0.7692	1.0000	0.0357	-0.7500	0.7335
6	0.3928	0.5555	1.0000	0.0370	-0.6071	0.5185

表 9: $N_{attack} = 4$ の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0370	0.8958	0.4667	0.3929	-0.4296	0.5029
4	0.2592	0.7777	0.9333	0.0741	-0.6740	0.7036
5	0.1666	0.7999	0.7143	0.1481	-0.5476	0.6518
6	0.5384	0.4762	0.9286	0.1154	-0.3901	0.3608

図 2: $N_{attack} = 1, 2, 4$ における AIRF ランキング

の配置数 N_{attack} を 1 とし, Fake コンテンツ数が $F = 1,024$, Bot の要求比率が $\rho = 0.1$ であるときに, Fake オリジンサーバと Bot の配置方式を変化させたときの予測精度を比較する.

配置方式がそれぞれ DCb 方式, DCf 方式, Random 方式のとき, Label1 となるノードはそれぞれ 8, 8, 10 箇所である. 教師データ数を 3 から 6 に変化させたときの, FNR, FPR, それぞれの差異を表 10, 11, 12 に示す.

表 10: DCb 方式の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.3796	0.3809	0.0000	0.5676	0.3796	-0.1866
4	0.5833	0.0555	0.8333	0.0833	-0.2499	-0.0277
5	0.3142	0.5000	0.0000	0.3889	0.3142	0.1111
6	0.4762	0.1333	0.2000	0.2571	0.2762	-0.1237

表 11: DCf 方式の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.3703	0.0476	1.0000	0.0278	-0.6296	0.0198
4	0.5000	0.0000	1.0000	0.0000	-0.5000	0.0000
5	0.2857	0.2777	0.6667	0.1714	-0.3809	0.1063
6	0.3904	0.1333	0.8333	0.0588	-0.4428	0.0745

表 12: Random 方式の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.2549	0.5185	0.2500	0.8000	0.0049	-0.2814
4	0.5000	0.4166	0.3750	0.2353	0.1250	0.1813
5	0.0707	0.9166	0.2857	0.4412	-0.2150	0.4754
6	0.3232	0.5714	0.4286	0.2121	-0.1054	0.3593

5.6 ネットワークトポロジによる差異

At Home Network と, Allegiance Telecom の 2 種類のトポロジにおける予測精度を比較する. DCb 方式で Fake オリジンサーバと Bot の配置数 N_{attack} を 1 とし, Fake コンテンツ数が $F = 1024$, Bot の要求比率を $\rho = 0.3$ とする.

At Home Network と Allegiance Telecom において, Label1 となるノードはどちらも 8 箇所である. 教師データ数を 3 から 6 に変化させたときの, FNR, FPR, それぞれの差異を表 13, 14 に示す. Allegiance Telecom では, 次数交錯方式における FNR の改善が大きく見られ, 無差別選択方式においても FPR の改善が大きく見られた. これは, 5.4 節においても指摘した, AIRF が原因だと考えられる. 図 3 は, 各トポロジの AIRF が大きいものに並べたものである. Allegiance Telecom の方が各ノードに注入された Fake コンテンツ数が大きくキャッシュヒット率に与える影響が増加することによって, より正確な特徴量行列が獲得できていると考えられる.

いずれの方式においても精度が良い, Allegiance Telecom の教師数 4 の場合の, 各方式の予測トポロジと正解ラベルのトポロジを図 4 に示す.

表 13: At Home Network の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.3426	0.1428	0.0000	0.5946	0.3426	-0.4517
4	0.5000	0.0000	0.6667	0.2500	-0.1667	-0.2500
5	0.3809	0.0555	0.2000	0.3889	0.1809	-0.3333
6	0.4095	0.0000	0.4000	0.3429	0.0095	-0.3429

表 14: Allegiance Telecom の GCN の予測精度

教師	無差別選択		次数交錯		差異	
	FPR	FNR	FPR	FNR	FPR	FNR
3	0.0465	0.4762	0.2857	0.3488	-0.2391	0.1274
4	0.0387	0.5555	0.2857	0.1667	-0.2469	0.3888
5	0.0555	0.5555	0.5714	0.0000	-0.5158	0.5555
6	0.0952	0.5333	1.0000	0.0000	-0.9047	0.5333

Average injection ratio of fake content ranking

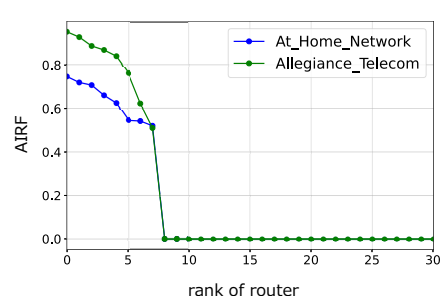


図 3: 各トポロジにおける AIRF ランキング

6. ま と め

本稿では, ICN の CPA に対する脆弱性を軽減するための新たなノード検出手法を提案した. 従来の CPA の防御法ではデジタル署名の利用が提唱されているが, 攻撃者が独自に作成した無意味なコンテンツをキャッシュさせる独自 Fake 型 CPA では, 署名が正当なため署名による防御が不可能であった. また, 独自 Fake 型 CPA はユーザからの要求が発生しないため, 検知も困難である. この問題に対し, 本稿では

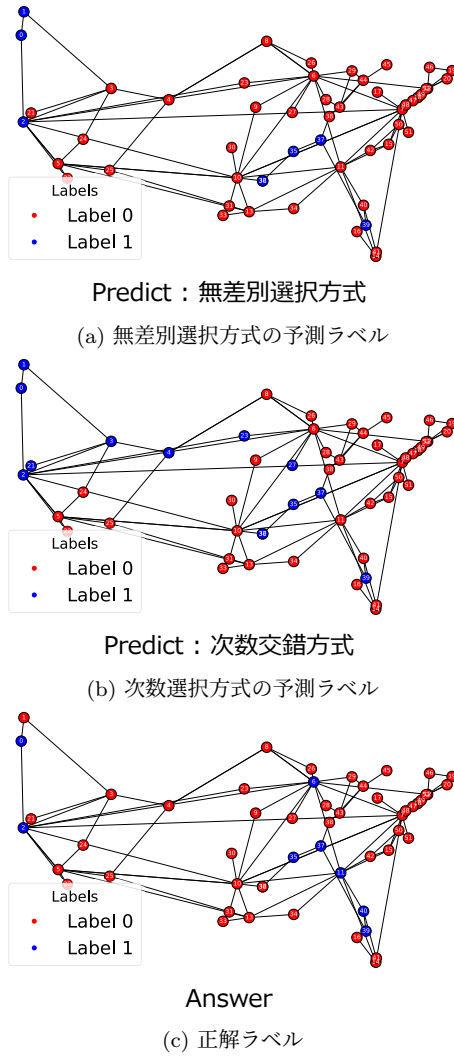


図 4: Allegiance Telecom, $N_{attack} = 1$, DCb 方式配置, 要求比率 $\rho = 0.3$, 教師数 4 の GCN による予測と正解ラベル

GCN を用いて Fake コンテンツが注入されやすい脆弱なノードを効率的に特定する手法を提案した。

提案手法では, ノードの AIRF をラベルとして利用し, 人気コンテンツのキャッシュヒット率を特徴量として学習を行った。この学習モデルにより, Fake 型 CPA に対する脆弱ノードを推測することが可能である。異なるトポロジを使用した評価や, 攻撃者の配置数などさまざまなパターンでの評価を行うことによって, GCN による予測の精度を評価した。本稿では, 無差別選択方式と次数交錯方式の 2 通りで, 予測精度の検証を行った。多くの場合で, 次数交錯方式は FNR を低減するのに有効であり, 無差別選択方式では FPR を低減するのに有効であるという評価が得られた。脆弱なノードの見逃しを防ぐためには, FNR が低減される次数交錯方式の方が優れていると言えるが, 一方で 5.4 節でも触れたとおり次数交錯方式では FPR が 1 になる場合も散見された。そのため, 今後は FPR が 1 にならないように FNR を低減できる教師データの与え方を考慮する必要があると考えられる。

また, 本稿の性能評価において教師データを増やしていく過程で精度の値の急激な変化が見られた。隠れ層の深さや特徴量の次元数など, より適切なモデルのハイパーパラメータを設定することによって, より高精度な GCN を構築していく必要があると考えられる。

謝辞 本研究成果は JSPS 科研費 23K21664, 23K21665, 23K28078 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] M. Zhang et al., A Survey of Caching Mechanisms in Information-Centric Networking, IEEE Communications Surveys & Tutorials, July 2015.
- [2] T. Nguyen, X. Marchal, G. Doyen, T. Cholez, and R. Cogranne, Content Poisoning in Named Data Networking: Comprehensive Characterization of real Deployment, IM2017, 2017.
- [3] N. Kamiyama, and T. Kudo, Investigating Impact of Fake-Type Content Poisoning Attack on NDN, GRASEC 2023.
- [4] W. Cui, et al., “Feedback-Based Content Poisoning Mitigation in Named Data Networking”, IEEE ISCC 2018.
- [5] T. Okada and N. Kamiyama, Name Management Using IOTA in ICN, IEEE DAG-DLT 2024, May 2024.
- [6] J. Zhou, et al., Graph neural networks: A review of methods and applications, AI Open, Volume 1, 2020, Pages 57-81.
- [7] Pujol-Perich et al., Unveiling the potential of Graph Neural Networks for robust Intrusion Detection, ACM SIGMETRICS Performance Evaluation Review, Volume 49, Issue 4, Pages 111-117.
- [8] L. Wu, P. Cui, J. Pei, and L. Zhao. Graph Neural Networks: Foundations, Frontiers, and Applications. Springer, Singapore, 725p.
- [9] Kipf, T. N, and Welling, M, Semi-supervised classification with graph convolutional networks. In International Conference on Learning Representations, 2017.