

面向QUIC的增量式网站指纹攻击方法

谭静文¹, 王焕然^{1*}, 韩 帅¹, 杨 武¹, 秦克伟²

(1. 哈尔滨工程大学计算机科学与技术学院, 黑龙江哈尔滨 150001; 2. 中移系统集成有限公司, 北京 100071)

摘 要: 快速用户数据报协议网络连接(Quick UDP Internet Connections, QUIC)旨在提供快速、稳定且安全的网络服务,其加密特性为加密恶意网站提供了保护伞。为了监管QUIC流量中加密恶意网站,网站指纹攻击成为研究热点。QUIC协议的渐进式部署特点,使得对新启用QUIC协议的网站实施早期WF攻击至关重要。攻击者必须持续监控网站是否开始支持QUIC协议,并在其启用QUIC协议后及时爬取少量QUIC流量,快速训练新的攻击模型。这种攻击可被称为增量式小样本网站指纹攻击。然而,训练样本的匮乏使得现有深度学习方法难以构建有效的深度特征。同时,现有方法均面向TCP协议提取初始特征,但这些初始特征无法表达QUIC协议特性,导致无法充分表征QUIC流量。这些问题导致现有方法在进行增量式小样本网站指纹攻击时性能下降。为解决此问题,本文分析了QUIC协议的特有特征及其对WF攻击的重要性。引入QUIC特有特征后,能更全面地表征网站的QUIC流量。同时,提出基于图论的增量式小样本网站指纹攻击方法,从全局视角挖掘潜在特征关联性,进一步提升攻击模型的表征能力。为模拟监测列表中新启用QUIC网站的出现情况,以一个月为间隔爬取两组网站流量数据。通过综合实验评估了所提特征与方法在新收集的QUIC网站流量上的有效性。相较于现有的先进方法,本文提出的方法在准确率上提升了1.01%~6.84%。

关键词: 加密恶意网站监管;QUIC协议;网站指纹攻击;增量识别;特征工程;图表示

基金项目: 国家自然科学基金(No.U22A2036, No.U21B2019, No.62272127, No.62572144);黑龙江省自然科学基金联合指导项目(No.LH2024F036);海南省自然科学基金高层次人才专项(No.622RC672)

中图分类号: TP393 **文献标识码:** A **文章编号:** 0372-2112(2026)04-1548-14

电子学报URL: <http://www.ejournal.org.cn>

DOI: 10.12263/DZXB.20251031

An Incremental Website Fingerprinting Attacks for QUIC Traffic

TAN Jingwen¹, WANG Huanran^{1*}, HAN Shuai¹, YANG Wu¹, QIN Kewei²

(1. School of Computer Science and Technology, Harbin Engineering University, Harbin, Heilongjiang 150001, China;

2. China Mobile System Integration Co., Ltd., Beijing 100071, China)

Abstract: Quick UDP internet connections (QUIC) protocol has been proposed to provide fast, stable and secure web services. Its encryption properties provide a protective shield for malicious websites. To detect malicious websites in QUIC traffic, website fingerprinting (WF) attacks have become a research hotspot. As the QUIC protocol is being gradually implemented, it is essential to conduct an early and rapid WF attack for new QUIC-enabled websites. Attackers must continuously monitor whether websites begin supporting the QUIC protocol, once QUIC is enabled, crawl a small amount of QUIC traffic for rapid training of new attack models. This type of attack may be termed an incremental few-shot website fingerprinting attack. However, the scarcity of training samples renders existing deep learning methods ineffective in constructing effective deep features. Moreover, current methods extract initial features from the transmission control protocol, which fail to capture the characteristics of the QUIC protocol and thus inadequately represent QUIC traffic. These limitations result in existing methods performing poorly when conducting incremental few-shot website fingerprinting attacks under QUIC traffic. To address the issues, we analyze QUIC-specific characteristics and their importance for WF attacks. The addition of QUIC-specific characteristics provides a more comprehensive representation of the QUIC-enabled websites. Meanwhile, a graph-based few-shot incremental website fingerprinting attack method is proposed to mine the potential characteristic relevance from a global perspective. It further improves the representation ability of the attack models. To simulate the emergence of new QUIC-enabled websites in the monitored list, we crawl the traffic of two groups of websites one-month apart. We conduct comprehensive experiments to evaluate the effectiveness of the proposed characteristics and method in the collected QUIC-enabled website traffic. Compared to the state-of-the-art method, the accuracy of the proposed method is improved by 1.01%~6.84%.

Keywords: regulation of encrypted malicious websites; QUIC protocol; website fingerprinting attack; incremental identification; feature engineering; graph-based representation structure

Foundation Item(s): National Natural Science Foundation of China (No.U22A2036, No.U21B2019, No.62272127, No.62572144); Natural Science Foundation Joint Guidance Project of Heilongjiang (No.LH2024F036); Natural Science Foundation High-level Talents of Hainan (No.622RC672)

0 引言

作为新型通用传输层网络协议,快速用户数据报协议网络连接(Quick UDP Internet Connections, QUIC)^[1]为用户提供了更快、更稳定且更安全的网络服务。相较于其他协议,QUIC 加密了更多信息,包括传输层安全性协议(Transport Layer Security, TLS)握手信息、流标识符等,使得加密通信中可用的明文极少^[2-3],有效保障了网络用户的隐私。然而,其加密特性为各类恶意网站提供了保护伞,使得监管人员无法有效监测并阻断其流量。网站指纹(Website Fingerprinting, WF)攻击可基于网站加密流量的元数据(大小、方向和时间)生成网站指纹^[4],以实现不解密情况下 QUIC 流量与特定网站的关联^[5],已逐渐成为网络安全领域的研究热点。

由于 QUIC 协议的部署是渐进且尚未完成的^[6],WF 攻击者监控列表中的部分网站可能不支持 QUIC 协议,无法预先爬取所有监控网站的 QUIC 流量。攻击者必须持续监控网站是否开始支持 QUIC 协议,并在其启用 QUIC 协议后及时爬取 QUIC 流量,以训练新的攻击模型^[7]。为实现对新启用 QUIC 协议网站的近乎即时的 WF 攻击,攻击者需要缩短爬取周期^[8]。这意味着攻击者仅能利用少量 QUIC 流量进行模型训练。此类 WF 攻击可称为增量式少样本网站指纹攻击(Incremental Few-Shot Website Fingerprinting attacks, IFSWF)。

现有主流的 IFSWF 攻击可分为两类:基于端到端深度学习^[9-12]和基于迁移学习^[8,13]的 IFSWF 攻击。基于端到端深度学习的 WF 攻击依赖于训练数据量,无法基于少量流量构建有效的特征表示。基于迁移学习的方法利用海量数据预训练强大模型,再用少量新任务数据对分类器进行微调。这类方法的有效性依赖于预训练数据的多样性。由于监控列表中初始支持 QUIC 的网站较少,其预训练数据的多样性有限。这导致预训练模型表征能力不足,进而限制了 IFSWF 攻击的有效性。此外,两类方法均沿袭了面向传输控制协议(Transmission Control Protocol, TCP)的 WF 攻击的特征^[12-14],未关注 QUIC 特有的属性特征,如 1-往返时间(1-Round-TripTime, 1-RTT)与 0-往返时间(0-Round-TripTime, 0-RTT)数据包的相关特征。这导致 QUIC 流量特征表征不完整^[15],限制了 IFSWF 攻击在 QUIC 协议下的性能。综上所述,现有方法的有效

性受限于对 QUIC 流量特征及深度特征的表征不足。

为提升表征能力,本文提出了 QUIC 特有特征以及一种新型 IFSWF 攻击模型。为更准确地表达 QUIC 流量,本文深入分析了 QUIC 协议,并探讨了 QUIC 特有特征与流量统计特征对 IFSWF 攻击的重要性。通过引入 QUIC 独有的特征,所提出的攻击模型实现了对 QUIC 协议更全面的表征。为构建有效的深度特征,面向 QUIC 协议提出基于图结构的网站指纹攻击方法(Graph-based WF for QUIC, GWF-QUIC)。该方法包含两个核心组件:基于图的特征表示(Graph-based Feature representation, GF-rep)与全局图特征聚合器(Global Feature aggregator, GF-agg)。GF-rep 是一种基于重要性权重的图结构,通过虚拟中心关联所有特征。GF-agg 利用图卷积网络(Graph Convolutional Networks, GCN)^[11]在重要性引导下聚合所有特征,从全局视角挖掘有用的深度特征,从而进一步提升攻击模型的表征能力。

本文的主要贡献如下。

(1) 提取了 43 个有利于分类的特征,包括 QUIC 独有的特征和统计特征。特征重要性分析结果表明,不同网站的 QUIC 的独有特征(如 1-RTT 和握手数据包信息)存在显著差异。

(2) 提出 GWF-QUIC 方法构建高效深度特征。该方法通过基于重要性的图结构关联所有物理特征,并基于全局特征聚合挖掘特征间潜在关联。依托全局深度特征,该方法显著提升了 QUIC 流量环境下 IFSWF 攻击的有效性。

(3) 基于收集的数据集,通过详细实验对提出的特征及 GWF-QUIC 方法进行实证评估。实验结果表明,相较于现有的先进方法,该方法在 QUIC 流量环境下的 IFSWF 攻击性能提升了 1.01% ~ 6.84%。

1 相关工作

1.1 基于传统深度学习的网站指纹攻击

由于大多数面向 QUIC 流量的 WF 攻击都源于面向 TCP 流量的 WF 攻击,首先回顾面向 TCP 的方法。Abe 等人^[16]首次将深度学习方法应用于 WF 攻击研究。他们采用堆叠自编码器(SAE)实现特征自动提取。与此同时,AWF^[17]评估了三种深度学习模型:堆叠去噪自编码器(Stacked Denoising Auto Encoder,

SDAE)、卷积神经网络(Convolutional Neural Network, CNN)及长短期记忆网络(Long Short-Term Memory, LSTM)。为从元数据序列中提取更有价值的特征,DF^[12]考虑了流量特性,并对基于CNN的模型进行了调整。为解决深度学习模型对网站流量变化敏感的问题,snWF^[18]、RF^[19]和Holmes^[20]被提出。其中,snWF基于快照集成技术降低神经网络的方差,RF通过信息泄露分析实现稳健的流量表征,Holmes通过分析网站流量的时空分布特性确保基于早期流量的稳健可靠检测。为了对抗用户的防御机制,Swallow^[21]通过深度学习模型提取流量中受防御机制影响较小的判别性特征,实现对匿名网络的高效WF攻击。考虑到用户会同时打开多个标签页导致网站识别不准确的问题,ARES^[22]被提出,将多标签攻击转化为多标签分类问题,从多个流量片段中提取局部特征来识别特定网站。为降低数据依赖性,Var-CNN^[9]设计了半自动特征提取机制,兼具人工特征提取与自动特征提取的优势,而GanDalf^[23]则利用半监督生成对抗网络生成更多的训练数据。由于上述方法所采用的特征均针对TCP协议设计,无法充分表达基于用户数据报协议(User Datagram Protocol, UDP)协议的QUIC流量。

随后,为面向QUIC流量实施WF攻击^[24],QUIC-CNN^[11]引入了新特征并优化了神经网络结构。为实现基于少量QUIC流量的高效攻击,WAPCN^[7]采用同一网站的TCP数据作为辅助。通过在特征空间中对齐同一网站的TCP与QUIC流量,该方法学习了网站感知型协议不变特征。由于这些方法未考虑QUIC特有的特征,其对QUIC流量的表达仍不完整。

此外,传统深度学习方法依赖海量数据训练高效特征提取器。由于IFSWF场景下支持QUIC的新网站样本量稀少,特征提取器的性能受到限制,难以适应当前QUIC渐进式部署下的IFSWF攻击需求。

1.2 基于迁移学习的网站指纹攻击

基于迁移学习的方法通过预训练构建强大的特征提取器,并允许分类器通过微调适应新任务^[25]。TF^[8]采用三重网络学习样本间的相似性。由于相似度度量可迁移至未见任务数据,CF^[26]通过改进三重损失函数并引入数据增强技术来提升迁移能力。由于基于TF的方法无法有效利用大规模辅助训练数据,TLFA^[13]应运而生。该方法利用交叉熵损失进行模型预训练,并将除softmax层外的所有层复用为特征提取器。为在小规模训练数据下进一步提升性能,MBL^[27]通过元学习理念优化目标任务。与此同时,TFAN^[28]在元学习过程中最小化查询特征图与支持特征之间的空间差异。考虑到历史数据复用可增强

预训练模型的表征能力,WFBDC^[29]运用迁移学习与多重相似度损失理念,缓解了历史数据集与当前数据集间的偏差。为实现攻击模型在新版本Tor及网站版本间的迁移,研究人员提出了CWFA^[30]与JAN^[31]。CWFA基于聚类技术将不同网站流量嵌入同一特征空间,JAN则通过缩小原始流量与新采集流量之间的特征空间差异及语义空间差异来实现迁移。遗憾的是,这些方法仍未能充分表征QUIC流量特性。此外,其有效性受限于预训练数据的多样性。由于监测列表中启用QUIC的网站数量有限,数据多样性严重不足,严重制约预训练特征提取器的效能,使其难以实现有效迁移。

为解决上述问题,本文分析了QUIC协议的独有特征,并提出了一种基于图的特征表示和特征聚合器。得益于对流量特征和深度特征的强大表示能力,所提方法显著提升了QUIC流量环境下IFSWF攻击的性能。

2 威胁模型

攻击场景如图1所示。攻击者是位于用户与网络服务提供商之间的被动监听者,他们能够捕获通信中的全部流量,但无法修改传输内容或解密数据包^[32]。攻击者的目标是从加密的QUIC流量中获取每个网站的独有特征,即指纹信息。随后攻击者利用这些指纹识别用户正在访问的网站。

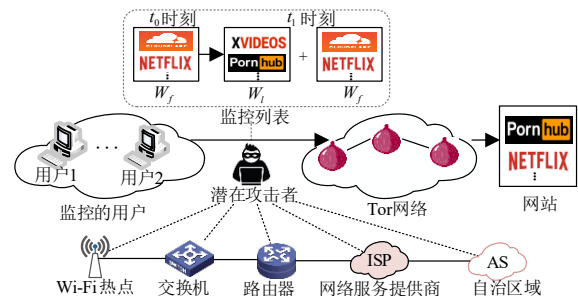


图1 威胁模型示意图

Figure 1 The threat model

一般来说,完整的网站指纹攻击过程分为数据采集、数据处理、攻击模型训练和攻击实施四个步骤。在数据采集阶段,攻击者模拟用户访问 m 个监控网站集 $W=\{w_0, w_1, \dots, w_m\}$ 的行为,并收集加密通信流量;在数据处理阶段,攻击者解析并提取不同网站的加密流量元数据(如数据包大小、方向等)^[33],形成流量痕迹traces;在攻击模型训练阶段,攻击者基于traces和其对应的网站标签训练分类器,获得攻击模型;在攻击实施阶段,攻击者监听用户的上网流量,并利用训练好的分类器发起攻击,推断用户当前产生的流量所属的监控网站。

在真实场景中,新启用 QUIC 协议的网站不断增多,监控网站集 W 将逐步扩展。具体而言,定义 $W_f = \{w_0, w_1, \dots, w_k\}$ 为 t_0 时刻已支持 QUIC 协议的旧网站, $W_l = \{w_{k+1}, w_{k+2}, \dots, w_m\}$ 定义为 t_1 时刻新启用 QUIC 协议的新网站。那么,随着时间 t 的推移,攻击者的监控网站集从 W_f 扩展为 $W_f \cup W_l$, 如图 2 所示。为实现对 W_l 实施近乎即时的识别,攻击者需要压缩上述四个攻击步骤的时间消耗。鉴于数据采集阶段是最耗时的,攻击者只能通过减少捕获的流量样本数量的方式缩减时间。因此,IFSWF 攻击者的目标是在 t_1 时刻,捕获较少的 W_l 的 QUIC 流量样本,更新攻击模型,实现面向 $W_f \cup W_l$ 的网站指纹攻击。



图2 监控网站集扩展示意图

Figure 2 The extensions for the monitored website set W

接下来,定义 QUIC 流量的特征集合 $F = \{f_0, f_1, \dots, f_n\}$, 其中 f_i 是第 i 个特征, n 是特征的个数。在此基础上,定义特征重要性集合 $\Theta = \{\theta_0, \theta_1, \dots, \theta_n\}$, 其中 θ_i 是特征 f_i 的重要性。根据特征重要性分析方法——条件熵、最小绝对收缩和选择算子 (Least Absolute Shrinkage and Selection Operator, LASSO), 特征重要性 Θ 可表示为条件熵 h 和 LASSO 线性系数 c 。对于每个 traces, 其在 W_f, W_l 中对应的网站类别用 y 表示。此外,定义 W_o 为用户在真实场景中可能访问的 W_f, W_l 以外的网页。

3 特征分析

网站中的部分资源或服务支持 QUIC 协议,而另一些仅支持 TCP,因此网站的访问流量中可能同时包含 QUIC 流量和 TCP 流量。在提取特征前,仅考虑 QUIC 流量并剔除 TCP 流量。鉴于部分流量统计特征仍适用于 QUIC^[34], 本文仍保留了这些特征。

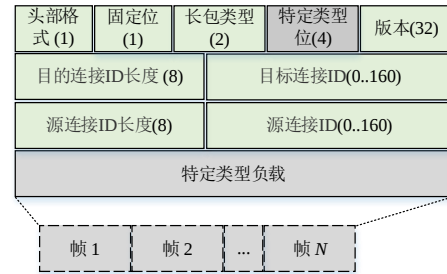
3.1 特征

3.1.1 QUIC 特有特征

与基于协议 TCP 的 TLS 流量不同,QUIC 依赖于 UDP 协议,后者不具备可靠的数据传输特性。为确保数据传输可靠性,QUIC 协议设计了若干特殊字段。这些字段会暴露数据传输模式的信息。

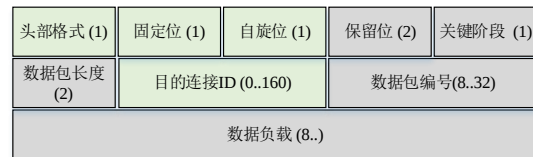
QUIC 协议在 RFC 9000^[35] 中实现了标准化。其结构包含报头包与帧,报头包可分为长报头与短报头两种格式,如图 3 所示。绿色字段表示未加密,灰色字段表示受加密或报头保护机制保护。通过可见字

段“头部格式”可识别报文头类型,“1”表示长报头,“0”表示短报头。



(a) 长报头

(a) Long header



(b) 短报头

(b) Short header

图3 QUIC 包头示意图

Figure 3 The header format of a QUIC packet

长报头用于连接建立前的数据传输,支持初始报文、握手报文、0-RTT 报文和重试报文这四类报文,可通过可见字段“长包类型”进行区分。初始报文与握手报文用于建立连接,0-RTT 报文与重试报文则用于在握手完成前传输“早期”数据。若客户端与服务器不是首次通信且短时间内通信过,采用 0-RTT 模式握手。相反地,则采用 1-RTT 握手模式。连接建立后,1-RTT 报文使用短报头作为实际数据传输的紧凑格式。

由于 QUIC 连接的建立依赖于初始报文和握手报文,本文采用初始数据包与握手数据包的数量及大小作为 QUIC 独有的特征。鉴于 1-RTT 报文承载实际数据,提取 1-RTT 数据包的大小相关特征。同时,由于 0-RTT 数据报文用于早期数据传输,也选取其大小相关特征。此外,本文发现 1-RTT 数据包尺寸不超过 1 250,故采用特定尺寸阈值进行特征提取。所选 QUIC 特有特征的符号与定义如表 1 所示。

3.1.2 流统计特征

流量统计特征已被证实能为基于 QUIC 的网站分类提供有益信息。流量中可获取的元数据包含每个数据包的大小、计数和时间。因此,从三个角度引入流量统计特征,所选流量统计特征的符号与定义如表 2 所示,具体说明如下。

(1) 与数据包大小相关的特征。这些特征通常揭示网站嵌入内容的规模,包括总数据包大小、整个访

表1 QUIC特有特征及对应的符号

Table 1 QUIC-specific features and corresponding symbols

No.	符号	定义
1	init_pkt_cnt	初始数据包数量
2	init_pkt_size_tot	初始数据包总大小
3	one_rtt_pkt_cnt	1-RTT数据包数量
4	one_rtt_pkt_size_tot	1-RTT数据包总大小
5	one_rtt_pkt_size_avg	1-RTT数据包平均大小
6	zero_rtt_pkt_cnt	0-RTT数据包数量
7	zero_rtt_pkt_size_tot	0-RTT数据包总大小
8	zero_rtt_pkt_size_avg	0-RTT数据包平均大小
9	handshake_pkt_cnt	握手数据包数量
10	handshake_pkt_size_tot	握手数据包总大小
11	certain_cnt_tot	大小为1 250的包数量

问的平均大小以及上行(或下行)流量的大小标准差。为分析数据包大小分布,选取上行(或下行)流量中小、中、大数据包的占比作为特征。由于前向数据包与链接建立及网站结构模板相关,本文同时提取前30个数据包的总大小作为特征。

(2)数据包计数相关的特征。此类特征反映网站嵌入内容的数量,包括所有访问及上行(或下行)流量每秒传输的总数据包计数。同理,提取交互次数作为特征。

(3)时间相关的特征。尽管时间特征受网络环境影响,但研究表明其具有积极作用^[9]。包括上行(或下行)方向相同数据包间的总持续时间与平均间隔时间,即上行(或下行)流量的数据包平均间隔、最大间隔与最小间隔。

3.2 特征重要性分析

为分析上述特征对WF攻击的重要性,将网站域名视为已知变量 Y ,每个特征视为未知变量 X 。由于提取的特征中同时存在线性相关特征与离散特征,本文采用最小绝对收缩和选择算子(LASSO)^[36]作为重要性分析方法。LASSO是一种回归分析方法,它通过自动应用L1正则化使部分特征权重归零,从而实现特征选择,其回归的损失函数可定义为

$$J(c) = \frac{1}{2n} (X_c - Y)^T (X_c - Y) + \alpha \|c\|_1 \quad (1)$$

其中, n 为样本数量, c 为待调优的常数系数, $\|c\|_1$ 表示L1范数, α 是L1正则化的超参数。当 $\alpha=0.0001$ 时,LASSO的线性系数 c 如图4(a)所示。

此外,本文试图探索WF攻击中每个特征的信息增益,因此额外选择条件熵^[37]作为分析方法。对于给定 X ,条件熵定义为 Y 在 X 上的条件概率分布熵的期望值,其计算方式为

表2 提取流统计特征及对应的特征符号

Table 2 Flow statistical features and corresponding symbols

No.	符号	定义
12	pkt_size_tot	总流量数据包大小
13	pkt_size_avg	平均流量数据包大小
14	up_pkt_size_tot	上行流量数据包大小
15	up_pkt_size_avg	平均上行数据包大小
16	up_pkt_size_std	上行数据包大小标准差
17	down_pkt_size_tot	下行流量数据包大小
18	down_pkt_size_avg	平均下行数据包大小
19	down_pkt_size_std	下行数据包大小标准差
20	top_30pkt_size_tot	前30个数据包总大小
21	pkt_cnt	流量数据包计数
22	up_pkt_cnt	上行总数据包计数
23	down_pkt_cnt	下行总数据包计数
24	pkt_size_s	每秒传输数据包大小
25	pkt_cnt_s	每秒传输数据包计数
26	up_pkt_cnt_s	每秒上行数据包计数
27	down_pkt_cnt_s	每秒下行数据包计数
28	round_cnt	交互次数
29	duration_time	流量持续时间
30	up_pkt_iat_tot	上行数据包总间隔时间
31	up_pkt_iat_avg	上行数据包的平均间隔时间
32	up_pkt_iat_max	上行数据包的最大间隔时间
33	up_pkt_iat_min	上行数据包的最小间隔时间
34	down_pkt_iat_tot	下行数据包总间隔时间
35	down_pkt_iat_avg	下行数据包的平均间隔时间
36	down_pkt_iat_max	下行数据包的最大间隔时间
37	down_pkt_iat_min	下行数据包的最小间隔时间
38	up_s_ratio	上行数据包里小型数据包所占比例
39	up_m_ratio	上行数据包里中型数据包所占比例
40	up_l_ratio	上行数据包里大型数据包所占比例
41	down_s_ratio	下行数据包里小型数据包所占比例
42	down_m_ratio	下行数据包里中型数据包所占比例
43	down_l_ratio	下行数据包里大型数据包所占比例

$$h(Y|X) = \sum_x p(x) h(Y|X=x) \\ = - \sum_{xy} p(x,y) \log p(y|x) \quad (2)$$

其中, p 表示概率分布。对于给定特征,条件熵值越高,其重要性越低。各特征条件熵值 h 如图4(b)所示。

在图4中,从暖色(如黄色)到冷色(如蓝色和紫色),特征的重要性逐渐增加。通过分析结果,可得出以下三点结论。

(1)QUIC协议的独有特征有利于分类。由图3可见,编号为4和5的特性具有较中等的 c 值和极低的 h 值。由于1-RTT主要用于传输数据,其大小代表

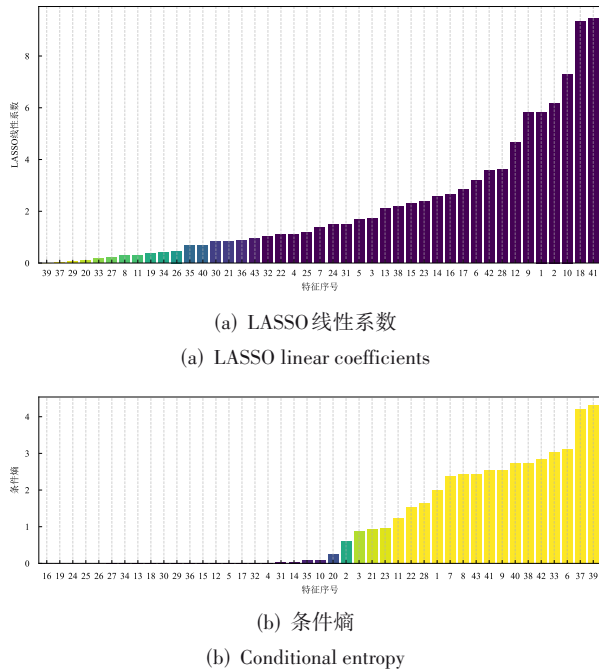


图4 特征重要性排序

Figure 4 Ranking of feature importance

网站包含的内容量,为WF攻击提供了有用信息。编号为2和10的特征具有较高的 c 值和较低的 h 值。由于所有资源提供者都需要握手过程,初始数据包和握手数据包的计数(或大小)隐含反映了支持QUIC协议的资源数量。

(2)数据包的大小和数量在QUIC协议中依然重

要。流量交互次数、总大小和平均大小是关键特征,这可以从编号为12~18的特征结果中看出。如前所述,网站提供的服务和显示的内容决定了传输数据的大小。

(3)时间相关特征提供了一定的信息价值。每秒传输的数据包数量与大小(特征编号24和25)获得中等偏下的 c 值和极低的 h 值。由于这些特征代表网站并行传输的数据量,受网络环境影响,并不稳定。

将特征的重要性定义为 $\Theta = \{\theta_0, \theta_1, \dots, \theta_n\}$,其中 θ_i 表示第 i 个特征的重要性。特征重要性 Θ 将成为GWF-QUIC方法的附加信息,帮助其更有效地学习网站指纹。GWF-QUIC方法能进一步挖掘未被发现的特征关联性,即使某些特征看似不重要,也依然保留。

4 方法

4.1 概述

为挖掘重要性各异的特征间潜在的全局关联性,本文提出基于图的WF攻击方法GWF-QUIC,其概述如图5所示。该方法由基于图的特征表示GF-rep和全局特征聚合器GF-agg组成。GF-rep采用带虚拟中心的加权图结构关联所有特征。在GF-rep中,节点代表各类特征,边权重是第3.2节生成的特征重要性值。GF-agg通过图神经网络(Graph Convolution Networks, GCN)从全局视角学习有利于分类的潜在特征。结合QUIC特有特征与全局特征聚合机制,该方法可生成稳定的网站指纹,显著提升QUIC流量环境下IFSWF攻击的性能。

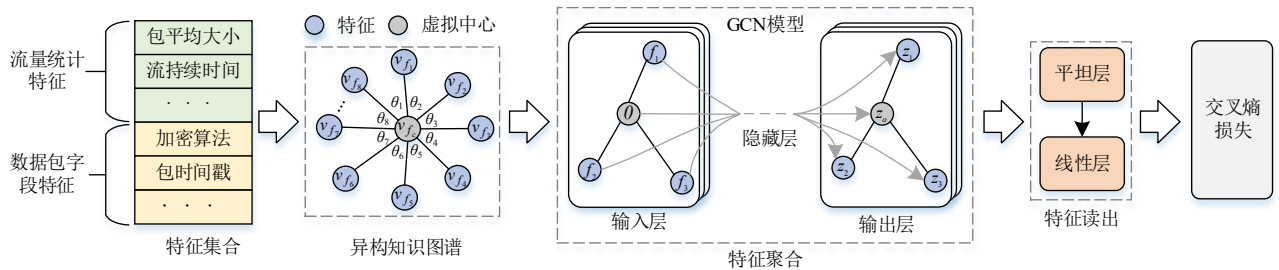


图5 GWF-QUIC方法示意图

Figure 5 Illustration of the GWF-QUIC method

4.2 基于图的特征表示

为关联线性相关或离散特征^[38-39],基于特征集 $F = \{f_0, f_1, \dots, f_n\}$ 构建图结构 $G = \{V, E\}$,其结构如图5所示。 $V = \{v_c, v_{f_0}, \dots, v_{f_n}\}$ 表示图 G 的节点集,其中 v_c 为虚拟中心节点, v_{f_i} 为特征 f_i 对应的节点。节点集 V 的大小为 $n+2$ 。 E 表示 G 的边集。每个节点 v_{f_i} 均与 v_c 相连。 v_c 节点不提供任何信息,仅用于关联所有特征。

对于给定的 G ,基于所有节点的特征,定义特征

矩阵为 $H = \{0, f_0, f_1, \dots, f_n\}$,其中 v_c 对应0。同时,基于节点之间的连接关系定义邻接矩阵 $A = (a_{ij})_{n \times n}$ 。其中,

$$a_{ij} = \begin{cases} 1, & \text{if } \langle v_i, v_j \rangle \in E \\ 0, & \text{else} \end{cases} \quad (3)$$

表示节点 v_i 与 v_j 间关系。通过这种方式,任意两个特征节点均可通过 v_c 实现互通。由于GF-rep将多个线性相关或离散的特征节点全部连接起来,后续基于图神经网络的GF-agg能够通过卷积操作聚合全局特

征,并学习特征之间的非线性依赖关系,从而提取有利于分类的高层深度特征。

为指导 GF-agg,引入第3节所述的特征重要性集 $\Theta=\{\theta_0, \theta_1, \dots, \theta_n\}$,用于计算 E 的权重。由于所有特征节点仅 v_c 相连, A 的计算公式可简化为 $a_{ci}=\theta_i$ 。

4.3 全局特征聚合

为挖掘特征的全局相关性,采用 GCN 作为攻击模型的框架^[40]。该模型专为图数据设计,能通过结构化的邻域聚合操作,自动地学习特征节点之间的高阶、非线性交互,提取全局潜在特征,从而学习稳定的潜在特征。

基于上述图 G 的邻接矩阵 A 与节点特征矩阵 H , GCN 执行如下卷积操作:

$$H^{(l+1)} = \sigma(\hat{D}^{-\frac{1}{2}} \hat{A} \hat{D}^{-\frac{1}{2}} H^{(l)} W^{(l)}) \quad (4)$$

其中, $\hat{A}=A+I$, I 为单位矩阵, D 为 $\hat{A}$ 的度矩阵, σ 为激活函数, W 为 GCN 的参数矩阵, l 为第 l 层。 D 为当前节点关联的总关系数,可通过累加矩阵 A 的每行计算得出。在无向图中, D 的计算方式如下:

$$D_{ij} = \begin{cases} \sum_p A_{ip}, & \text{if } i=j \\ 0, & \text{else} \end{cases} \quad (5)$$

经过 l 轮消息传播后,节点 v_{f_l} 的特征可定义如下:

$$H_{v_{f_l}}^{(l+1)} = \sigma \left(\sum_{v_{f_l} \in \text{NBR}_{v_{f_l}}} \hat{D}^{-\frac{1}{2}} \hat{A} \hat{D}^{-\frac{1}{2}} H_{v_{f_l}}^{(l)} W^{(l)} \right) \quad (6)$$

其中, $\text{NBR}_{v_{f_l}}$ 是节点 v_{f_l} 的邻居节点集合。由于 NBR_{v_c} 包含所有 v_{f_l} 节点,首轮消息传播使 v_c 节点获取 F 节点的聚合特征。在第二轮消息传播中,聚合特征从 v_c 传递至所有 v_{f_l} 节点。由此,仅需两层卷积即可实现全局消息传播,这有助于挖掘特征间的潜在关联性。

GCN 基于上述公式为每个节点生成嵌入向量。通过基于展平读取层的操作,可从所有节点嵌入中聚合出统一的图嵌入 emb_g 。最终, GWF-QUIC 方法以增加不同网站图嵌入向量差异为目标训练 GCN。基于交叉熵函数的损失函数定义为

$$L = -\frac{1}{N} \sum_{i=1}^N y_i \log(\text{softmax}(\text{emb}_g)) \quad (7)$$

其中, N 为网站流量样本数量, y_i 为对应网站流量第 i 个图样本的标签。通过最小化损失函数 L , 可将每个图分类至其所属类别。

与传统依赖复杂的人工特征交互的机器学习方法不同^[41], GWF-QUIC 方法通过图结构自动挖掘简单特征间复杂的全局依赖关系。同时,与受限局部卷积感知的基于 CNN 的深度学习方法不同^[40], GWF-QUIC

方法通过图上的信息传播机制,实现了特征间的全局学习。总而言之, GWF-QUIC 方法实现了机器学习与深度学习的有效结合,引入了机器学习方法提取的特征重要性,并基于全局图结构进一步挖掘全局特征关联性,使得模型能够基于少量训练数据学习到稳定、有效的分类特征。

5 实验评估

5.1 数据采集和处理

5.1.1 数据采集

为收集 QUIC 流量,使用位于美国硅谷的虚拟云机器模拟访问网站的行为。该机器搭载第 12 代英特尔®酷睿™i5-12400 处理器(2.5 GHz),配备 32 GB 内存,并安装 Windows 2022 服务器系统。选取 20 个主流网站作为爬取目标,通过 Selenium(版本 3.141)驱动 Chrome 浏览器(版本 122.0.6261.95)自动执行模拟操作,并基于 tshark 工具进行流量捕获。其中, W_f 的 10 个网站的流量采集于 2024 年 3 月,用于模拟原始监控列表; W_l 的 10 个网站的流量采集于 2024 年 4 月,用于模拟需新增至监测列表的新启用 QUIC 协议的新网站。模拟过程中, Chrome 浏览器始终处于启用 QUIC 协议的模式。爬取过程中,每个网站的最大允许加载时间为 45 s,加载失败的样本不予保留。为确保流量在页面渲染前完成传输,在加载成功后需要额外等待 5 s。此外,每个网站每次访问 50 次后需间隔 5 h,以避免 IP 封禁。最终,特定网站产生的全部流量将保存为 PCAP 文件。为确保数据的平衡性,从每个网站选取 260 个样本作为 traces 样本,最终构成封闭世界数据集。

此外, 2026 年 1 月,按照 VisQUIC^[42] 中提及的 URL 列表,额外采集了 9 000 个 URL 页面的 QUIC 流量 W_o ,用于模拟真实开放场景中用户可能访问的其他网页。由于已公开的启用 QUIC 协议的网站有限, W_o 与监控网站可能归属同一网站,但其访问的 URL 页面并不同。每个 URL 仅访问一次,共 9 000 个 traces 样本,最终构成开放世界数据集。

5.1.2 数据处理

首先,针对所有 QUIC 数据包,依据报文的第 1 位区分长报头与短报头,并根据长报头中的第 3~4 位区分握手数据包、初始数据包和 0-RTT 数据包。其次,记录所有类型数据包的大小、方向及到达时间。再次,基于大小、方向和到达时间生成表 1 和表 2 中的 43 个特征。最后,所有特征经归一化处理后构造为一维向量 $F'=\{f_0, f_1, \dots, f_{42}\}$ 。网站中的部分资源或服务支持 QUIC 协议,而另一些仅支持 TCP,因此网站的访问流量中可能同时包含 QUIC 流量和 TCP 流量。此

时需要剔除基于 TCP 的 TLS 流量,仅保留 QUIC 流量。

5.2 对比方法和实验设置

5.2.1 对比方法

为验证 GWF-QUIC 方法的有效性,选取 8 种最先进的方法作为基准。

(1) CUMUL^[14]。该方法是 TCP 环境下最优的机器学习型 WF 攻击方案,采样了 104 个特征(包括数据包大小的累积值与方向)。

(2) DF^[12]。这是基于深度学习的经典 WF 攻击方法。该方法面向 TCP 流量,根据流量方向调整基于卷积神经网络的模型,其准确率超过 98%。

(3) Var-CNN^[9]。该方法采用半自动特征提取机制,结合人工特征提取与自动特征提取,有效降低了网站指纹攻击对训练数据量的依赖性。

(4) QUIC-CNN^[11]。该方法将卷积神经网络引入 QUIC 场景中,采用数据包大小与方向作为特征。

(5) MM-CNN^[10]。该方法采用多模态卷积神经网络对 QUIC 流量中的 Web 服务进行分类,特征包含数据包元数据序列及流量统计特性。

(6) TF^[8]。这是一种经典的基于迁移学习的小样本 WF 攻击方法,通过学习适合 traces 的度量指标,实现对新网站的快速适应。

(7) TLFA^[13]。该方法利用海量辅助数据训练任务无关的嵌入模型,再通过少量任务特定数据对分类器进行微调。其中,支持向量机(Support Vector Machines, SVM)被设定为 TLFA 的分类器。

(8) WAPCN^[7]。该方法采用特征对齐理念,以缩小 TCP 流量与 QUIC 流量间的特征分布差异为目标训练攻击模型,使其能够基于更少的 QUIC 样本,提升新启用 QUIC 网站的分类性能。

其中, CUMUL 为基于机器学习的 WF 攻击方法,侧重 TCP 协议的特征选择工作。DF 和 Var-CNN 是面向 TCP 协议的传统深度学习 WF 攻击方法。这些基线用于验证 TCP 流量特征对 QUIC 流量的适用性; QUIC-CNN 和 MM-CNN 是面向 QUIC 协议的攻击方法,深入分析了 QUIC 流量特征,故被选作基线; TF、TLFA、WAPCN 是基于迁移学习的小样本网站指纹攻击,用于对比验证增量小样本 WF 攻击场景的性能。

5.2.2 实验设置

本文分别采用条件熵 h 和 LASSO 线性系数 c 作为特征重要性 θ 。因此, GWF-QUIC 方法具有两种形式: GWF-QUIC- h (将 h 作为 GF-rep 中边权重) 和 GWF-QUIC- c (将 c 作为 GF-rep 中边权重)。GF-agg 模型由两个图卷积层和两个线性层组成,分别包含 8、16、64 和 10 个隐藏神经元。每个图卷积层后紧接 1 个批量归一化层和 1 个概率为 0.1 的 dropout 层。训练周期设

为 200 次, 学习率 β 设为 0.001, 采用 Adamx 优化算法。

(1) 封闭世界场景设置。封闭世界场景假设用户仅访问监控网站集中的网站, 用于评估攻击模型能否准确地将 traces 分类到指定监控网站。由于目标是在增量式少样本设置中评估 WF 攻击性能, W_f 的训练数据充足而 W_l 的训练数据有限。对于 W_f , 从每个网站选取 208 个训练样本和 52 个测试样本。对于 W_l , 从每个网站选取 2、6、10、14、18 个训练样本记为 train_size, 并保留 52 个测试样本。对于 DF、QUIC-CNN、MM-CNN、Var-CNN 等端到端学习的对比方法, 在 t_0 时刻使用 W_f 全部训练样本训练分类器, 在 t_1 时刻使用 $W_f \cup W_l$ 全部训练样本训练分类器; 对于 TF、TLFA、WAPCN 三种基于迁移学习的对比方法, 在 t_0 时刻使用 W_f 全部训练样本预训练一个表示学习模型。随后, 在 t_1 时刻利用表示学习模型生成 $W_f \cup W_l$ 训练样本的嵌入向量, 并基于嵌入向量微调一个分类器。最终测试阶段, 所有方法使用 $W_f \cup W_l$ 全部测试样本进行评估。基于该设置, 评估过程可遵循增量式少样本网站指纹攻击的场景。由于封闭世界的目标可视为多分类任务, 采用分类指标准确率 (Accuracy)、宏精度 (Macro-precision)、宏召回率 (Macro-recall) 及宏 F_1 分数 (Macro- F_1 -score) 评估效能。

(2) 开放世界场景设置。开放世界场景假设用户可访问更多网站, 即使这些网站不在监控网站集中。该场景用于评估攻击模型能否区分监控网站与非监控网站。在训练过程中, 从 W_o 的全部数据中随机选择与监控网站训练数据量等量的 traces, 作为额外类别参与分类器训练。在测试过程中, 除封闭世界的测试样本外, 额外从 W_o 剩余的样本中随机选择 1 000 ~ 6 000 个 traces 进行测试。开放世界设置侧重于评估受监控网站与未监控网站的分类能力, 可视为二元分类任务。同时, 考虑到开放世界的 URL 规模远大于封闭世界监控网站样本, 本文使用平均精度 (Average Precision, AP) 以求更准确地反映模型在类不平衡场景下的表现。

5.3 封闭世界有效性评估

为验证 GWF-QUIC 方法的性能, 开展了封闭世界有效性实验, 实验结果如图 6 所示。由图 6 可知, 随着训练样本 train_size 值降低, 所有方法的性能均出现衰退。本文所提出的 GWF-QUIC 在所有 train_size 设置的准确率、宏精度和宏 F_1 分数均显著优于所有基线方案, 比 CUMUL 方法提升了 1.01% ~ 6.84%。而 GWF-QUIC 在 2-shot 时宏召回率不是最优, 这是因为训练样本极少时, GCN 模型容易过拟合。尽管如此, GWF-QUIC 方法依旧是相对最优的方法, 仅需 14 个训练样本即可在所有指标上达到 90% 以上的性能, 展现出极强的数据稀缺场景下的鲁棒性。GWF-QUIC- h

与 GWF-QUIC-c 具有相似的性能趋势,但 GWF-QUIC-h 在所有指标上均以 0.2% 的优势超越 GWF-QUIC-c。性能差异归因于条件熵所依赖的非线性概率建模,其表达能力显著强于 LASSO 的线性特征选择机制,因此 GWF-QUIC-h 展现出更优的效果。而 CUMUL 缺乏协议特异性特征导致整体性能难以提升。作为深度学习学习方法,DF、QUIC-CNN 和 MM-CNN 在数据不足时难以挖掘有效特征。DF 结构深且复杂,故更易受数据量影响。尽管 TF 和 TLFA、Var-CNN 通过迁移学习或者半自动化特征挖掘缓解了数据依赖性,但其性能仍

受限于对 QUIC 流量特征的不完整表征。由于没有大量 TCP 流量辅助模型进行训练,WAPCN 方法性能最差且最不稳定。此外,CNN 模型仅能学习局部特征,导致上述基于 CNN 的方法无法挖掘输入特征间的潜在关联性。相较之下,GWF-QUIC 方法将机器学习方法提取的特征重要性融入 GF-rep 模型。这些额外信息有助于模型在少量训练数据的基础上学习更稳健的表征。同时,从 GF-agg 挖掘的潜在全局特征关联性有利于提取稳定的潜在特征。得益于机器学习与深度学习的有效结合,GWF-QUIC 方法的鲁棒性得到提升。

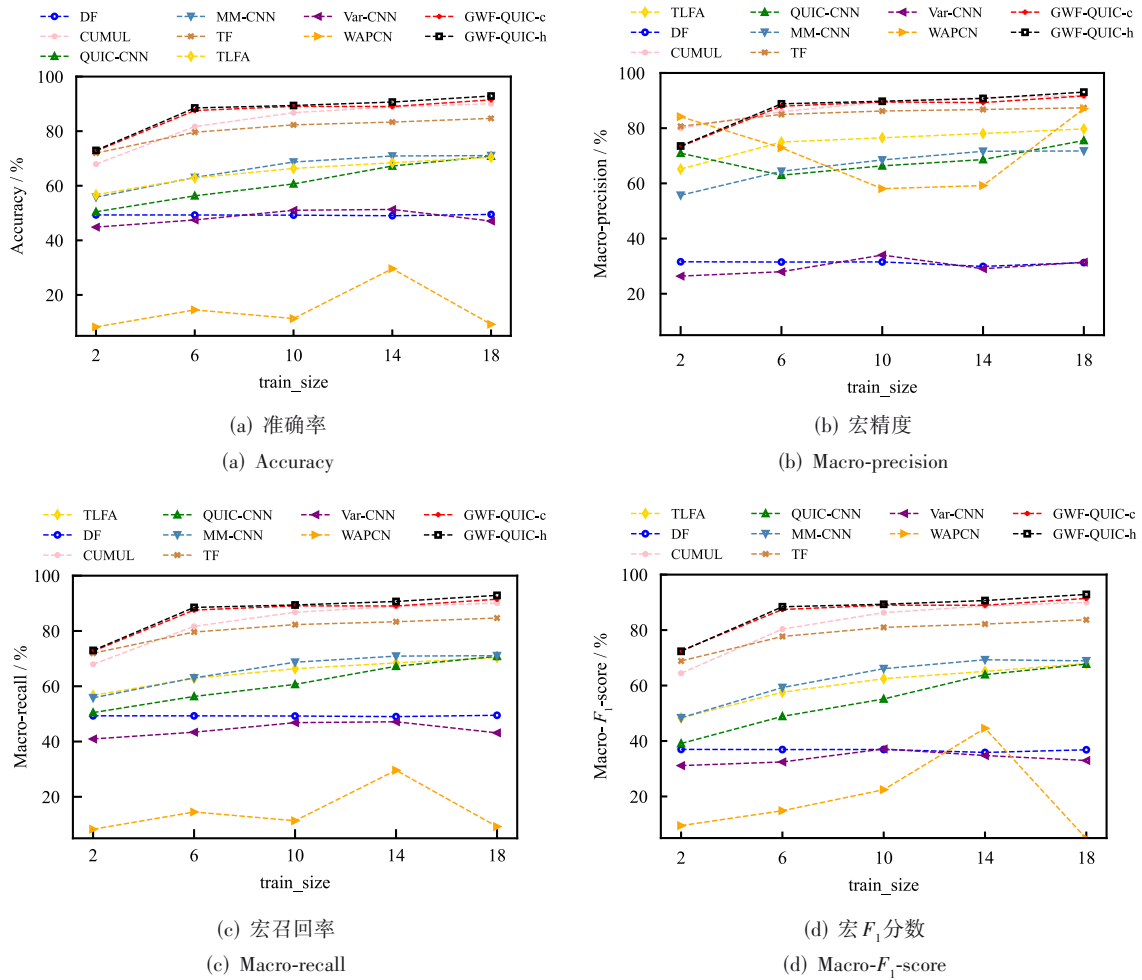


图6 有效性实验结果

Figure 6 Results of the validity experiment

5.4 开放世界鲁棒性评估

为评估方法在开放世界场景下的性能,该场景允许用户访问未被监控的网站。在训练过程中,为监控网站 W_f 和 W_l 中的每个网站各选择 18 个训练样本,共计 360 个监控网站样本。为保证监控网站和非监控网站的样本数相等,从非监控网站 W_o 中选择 360 个页面的 traces 作为训练样本。随后,在测试时,依旧

从每个监控网站 traces 中随机抽取 52 个测试样本;从非监控网站 traces 中分别抽取 1 000、2 000、3 000、4 000、5 000 和 6 000 个样本进行测试,用以模拟不同规模的开放世界。二分类的平均精度结果如图 7 所示。

由图 7 可知,在所有开放世界规模下,CUMUL、DF、MM-CNN、TF、TLFA、GWF-QUIC 方法的平均精度

极高, GWF-QUIC-h 最高。得益于 QUIC 独有的特性以及全局特征相关性挖掘, GWF-QUIC 方法在开放世界中表现极佳。CUMUL 通过人工设计的复杂特征取得了较好的性能, 但因缺乏 QUIC 特有特征, 性能难以持续提升。DF、TF 与 TLFA 采用的深度学习模型虽然能支持监控与非监控网站的二分类任务, 但由于所用流量特征过于单一, 性能提升空间有限。类似地, MM-CNN 虽分析了 QUIC 流量特征, 但模型在全局特征关联挖掘方面能力不足, 限制了其性能表现。Var-CNN 效果处于中等水平, 这归因于其所使用的时间特征较为原始, 未经过统计处理, 因而易受网络环境波动影响。相比于封闭世界, WAPCN 在开放世界中的性能有所提升, 这是因为该模型在训练过程中通过最小化同类样本的最大均值差异, 缩小了监控网站的内部差异, 从而扩大了监控网站与非监控网站之间的区分度。QUIC-CNN 性能较差, 尽管采用了与 DF、TLFA 相同的数据包大小和方向特征, 但其模型架构过于简单, 难以在海量非监控 URL 中形成清晰的分类边界。此外, 所有方法的平均精度均随开放世界规模的扩大而下降, 而 GWF-QUIC 的性能衰减最为缓慢, 表明其 CPFA 方法在更大规模的开放世界场景下具有更强的鲁棒性。

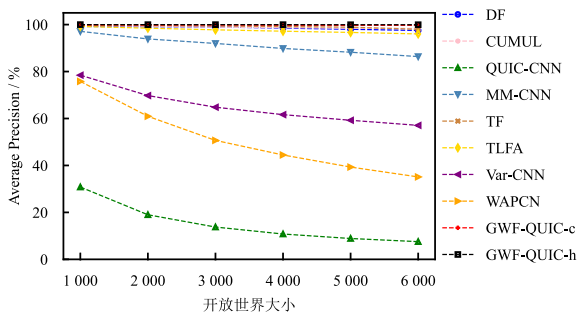


图7 开放世界鲁棒性实验结果

Figure 7 Results of the open-world robustness experiment

5.5 防御场景下鲁棒性评估

在实际场景中, 用户可在客户端部署防御策略, 通过对抗性攻击规避攻击者的检测^[43]。因此, 选取 FRONT^[44] 和 Tamaraw^[45] 两种典型防御策略评估所有方法在对抗性防御场景下的鲁棒性。其中, FRONT 方法通过注入双向随机干扰数据包实现对早期流量的混淆, Tamaraw 则通过人为延迟和恒定速率的固定大小数据包注入实现流量干扰。在实验过程中, 所有方法使用干净数据训练分类器, 使用防御后数据进行测试。最终结果如表3和表4所示。

由表3和表4可知, 当用户部署防御策略时, 所有方法均失效, DF、TF、TLFA 性能反而比 GWF-QUIC 方法性能有所提升。这是因为 DF、TF、TLFA 方法更侧重于从元数据序列中提取局部模式, 其 CNN 模型架

表3 FRONT防御场景下的各方法识别准确率 单位: %

Table 3 Accuracy of each method in the FRONT defense scenario

unit: %

train_size	2	6	10	14	18
DF	31.34	29.03	29.42	29.03	32.98
CUMUL	6.63	4.71	4.71	4.71	4.71
QUIC-CNN	5.0	5.0	5.0	5.0	5.0
MM-CNN	10.38	8.36	8.07	8.75	8.94
TF	23.84	28.65	30.58	31.63	32.01
TLFA	22.69	22.01	23.26	23.65	25.28
Var-CNN	20.0	26.53	20.96	25.86	24.03
WAPCN	11.63	8.75	13.07	6.92	10.48
GWF-QUIC-c	5.0	8.57	6.94	6.94	6.94
GWF-QUIC-h	5.0	5.97	6.93	6.93	6.93

表4 Tamaraw 防御场景下的各方法识别准确率 单位: %

Table 4 Accuracy of each method in the Tamaraw defense scenario

unit: %

train_size	2	6	10	14	18
DF	5.67	5.0	5.0	5.09	5.0
CUMUL	5.0	5.0	5.0	5.0	5.0
QUIC-CNN	7.88	6.05	5.96	7.02	6.06
MM-CNN	7.21	5.19	8.17	5.09	5.19
TF	4.61	4.61	4.61	5.19	5.67
TLFA	5.28	5.38	5.38	5.38	5.29
Var-CNN	5.0	5.0	5.0	5.0	4.9
WAPCN	5.0	5.76	5.76	4.9	5.57
GWF-QUIC-c	5.0	4.91	4.91	4.91	4.91
GWF-QUIC-h	5.01	5.01	5.01	4.91	5.01

构天然对局部特征位移具有鲁棒性, 更能适应防御造成的序列干扰。相对地, GWF-QUIC 的性能下降显著, 准确率在 5%~8% 之间。GWF-QUIC 方法的设计初衷是充分分析 QUIC 协议特征, 以实现更精准的 QUIC 网站指纹识别, 是一个分析性的攻击方法, 而非一个旨在对抗专门防御措施的鲁棒性攻击方法。因此, 当面对专门通过注入大量虚假数据包来干扰和污染流量统计特征的主动防御时, GWF-QUIC 所依赖的全局特征分布被改变, 这导致了其性能下降。特征易被干扰的特性反向印证了 GWF-QUIC 方法成功设计并捕捉到了对 QUIC 网站分类极为关键的特征。

5.6 消融实验

为评估本文的贡献, 分别进行特征贡献分析和方法贡献分析。

5.6.1 特征贡献分析

将 CUMUL、DF、QUIC-CNN、TF 和 TLFA 中使用的特征替换为本文提出的特征。变体版本分别记为 CUMUL-pf、DF-pf、QUIC-CNN-pf、TF-pf、TLFA-pf 和

WAPCN-pf。由于MM-CNN和Var-CNN的结构与本文提出的特征不匹配,没有对应变体。通过计算这些变体相对于原始版本的性能,可以评估所提特征的贡献。图8展示了各变体的评估结果。

对比图8和图6可知,CUMUL-pf、TLFA-pf和

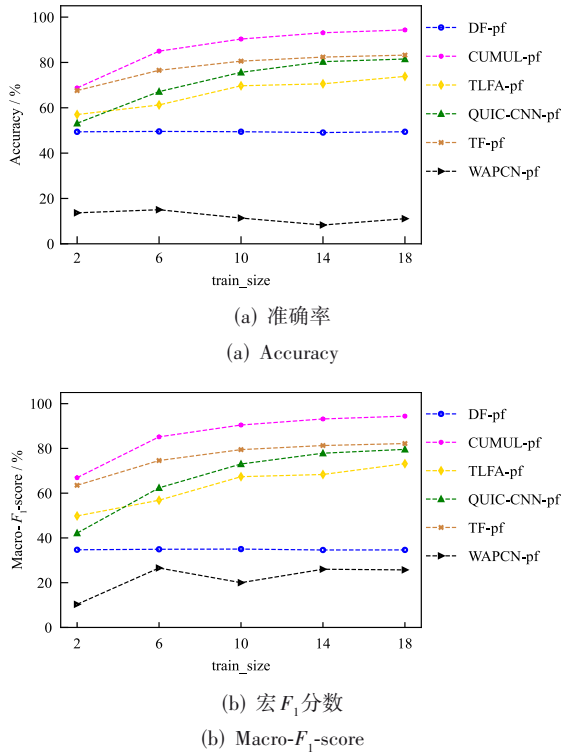


图8 全局特征贡献结果

Figure 8 Results of global feature contributions

QUIC-CNN-pf在所有场景中的性能分别提升了2.468%~4.45%、0.75%~5.52%和3.5%~17.87%。性能提升表明,本文提出的特征为QUIC流量下的IFSWF攻击提供了更多信息。相比之下,TF-pf、DF-pf和WAPCN-pf所有指标变化不大或稍微下降。由于这些方法的分类器是针对基于数据包元数据特征序列(数据包大小、方向等)的输入量身定制的,它们无法有效地从QUIC特有特征和统计特征中提取潜在特征。

为进一步量化QUIC特有特征的贡献,提出以下两种变体。

(1) GWF-QUIC-quic。这是GWF-QUIC方法的变体,仅使用QUIC特有特征训练分类器。该变体用于评估流量统计特征是否适用于面向QUIC的WF攻击。

(2) GWF-QUIC-flow。这是GWF-QUIC方法的一种变体,仅使用流量统计特征训练分类器。由于该版本去除了QUIC特有特征,可用于评估QUIC特有特征对WF攻击的贡献。

表5展示了基于GWF-QUIC且训练数据较少的变体所获得的结果。如表5所示,基于GWF-QUIC-flow的方法随着训练数据的增加而性能提升,而基于GWF-QUIC-quic的方法则呈现波动趋势。统计特征捕捉了网站流量的全局信息。随着样本数量的增加,统计信息趋于稳定状态分布,从而使性能逐步提升。由于QUIC同时具备0-RTT和1-RTT通信模式,其独有的特性遵循两种分布规律,这影响了基于GWF-QUIC-quic方法的稳定性。尽管如此,该方法仍保持着70%以上的准确率,最高达87%。这证明了QUIC的独有特性对IFSWF攻击的重要性。

表5 细粒度特征贡献分析结果

单位: %

Table 5 Results of the fine-grained analysis of feature contribution

unit: %

train_size	2		6		10		14		18	
指标	准确率	宏 F_1 分数	准确率	宏 F_1 分数	准确率	宏 F_1 分数	准确率	宏 F_1 分数	准确率	宏 F_1 分数
GWF-QUIC-quic	71.15	71.09	87.21	87.13	81.73	81.83	76.63	76.47	82.30	82.73
GWF-QUIC-flow	63.53	63.32	79.53	79.50	82.67	82.67	83.67	83.60	87.03	86.96
GWF-QUIC-h	72.96	72.35	88.51	88.40	89.40	89.30	90.67	90.64	92.88	92.85

5.6.2 方法贡献分析

为评估GWF-QUIC方法的贡献,提出了以下两种变体类型。

(1) MLP。鉴于MLP能够挖掘特征的全局相关性,提出一种变体方法来评估GWF-QUIC方法的贡献。该变体使用一维数组替代GF-rep,并将GF-agg替换为MLP。特征及其重要性沿行方向进行拼接。对应特征重要度 c 和 h ,变体方案分别标记为MLP-c和MLP-h。

(2) GWF-QUIC-n。这是GWF-QUIC方法的无权重变体,旨在评估该方法中特征重要度权重的贡献。

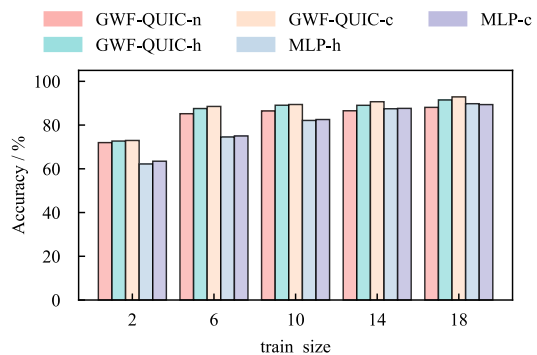
随着训练样本数量的增加,所有指标的结果如

图9所示。

在使用18个训练样本的结果中,GWF-QUIC方法比MLP-c高出3.61%。随着训练样本数量的减少,性能提升幅度逐渐增大。这证实了GWF-QUIC方法在学习全局特征方面优于MLP-c。由于GF-rep关联所有特征,且GF-agg执行了消息传播,GWF-QUIC能基于少量数据挖掘所有特征间的关联性。在使用2个训练样本的结果中,GWF-QUIC相较GWF-QUIC-n提升了约1.18%的性能。随着训练样本数量增加,提升幅度持续扩大,在使用18个训练样本时达到4.85%。这表明GWF-QUIC方法有效利用了预分析结果,成功

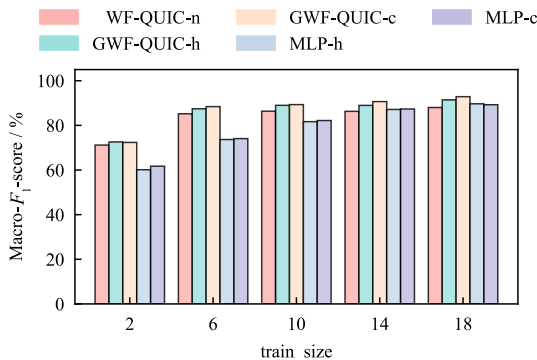
实现了机器学习与深度学习的融合。

综上所述,所有结果均表明本文所提出的新特征的贡献最大。在所有训练样本量下,新特征均带来 0.75% ~ 17.89% 的性能提升。由于 QUIC 特有特征能更精准地表征支持 QUIC 的网站,指纹构建仅需少量数据即可高效完成。GWF-QUIC 方法与特征重要性的贡献度相当,在 1.18% ~ 4.85% 之间。GWF-QUIC 方法能在全局图结构和特征重要性指导下挖掘特征关联性,仅需少量训练数据即可快速学习稳定的潜在特征。



(a) 准确率

(a) Accuracy



(b) 宏 F_1 分数

(b) Macro- F_1 -score

图9 GWF-QUIC方法的贡献结果

Figure 9 Contribution of the GWF-QUIC method

6 结束语

为有效监管网站的加密 QUIC 流量,本文设计了一种基于图的新型增量式少样本网站指纹攻击方法,即 GWF-QUIC。为更准确地表征启用 QUIC 的网站流量,分析了 QUIC 流量的独有特征,并探索了特征的重要性。GWF-QUIC 方法构建加权图,基于特征重要性建立特征关联,并通过全局特征聚合器挖掘全局特征相关性。依托精细化的 QUIC 独有特征,该方法获得更强的表征能力。此外,图结构将机器学习的特征分

析与深度学习的模型训练相衔接,助力 GWF-QUIC 方法实现更高效的新 QUIC 网站的早期识别。

参考文献

- [1] Langley A, Riddoch A, Wilk A, et al. The QUIC transport protocol: Design and internet-scale deployment[C]//Proceedings of the Conference of the ACM Special Interest Group on Data Communication. New York: ACM, 2017: 183-196.
- [2] Basyoni L, Erbad A, Alsabah M, et al. QuicTor: Enhancing tor for real-time communication using QUIC transport protocol[J]. IEEE Access, 2021, 9: 28769-28784.
- [3] Kumar P, Dezfouli B. Implementation and analysis of QUIC for MQTT[J]. Computer Networks, 2019, 150: 28-45.
- [4] Shreedhar T, Panda R, Podanev S, et al. Evaluating QUIC performance over web, cloud storage, and video workloads[J]. IEEE Transactions on Network and Service Management, 2022, 19(2): 1366-1381.
- [5] Gui Xiaolin, Cao Yuanlong, Huang Longjun, et al. A survey of QUIC-based network traffic identification[C]//Proceedings of the 12th International Conference on Mobile Networks and Management. Cham: Springer, 2023: 365-372.
- [6] Wang Tao, Goldberg I. Improved website fingerprinting on Tor[C]//Proceedings of the 12th ACM Workshop on Privacy in the Electronic Society. New York: ACM, 2013: 201-212.
- [7] Zhan Mengqi, Li Yang, Zhu Yongchun, et al. Website-aware protocol confusion network for emergent HTTP/3 website fingerprinting[J]. IEEE Transactions on Information Forensics and Security, 2023, 18: 2427-2439.
- [8] Sirinam P, Mathews N, Rahman M S, et al. Triplet fingerprinting: More practical and portable website fingerprinting with n-shot learning[C]//Proceedings of the 2019 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2019: 1131-1148.
- [9] Bhat S, Lu D, Kwon A, et al. Var-CNN: A data-efficient website fingerprinting attack based on deep learning[J]. Proceedings on Privacy Enhancing Technologies, 2019, 2019(4): 292-310.
- [10] Luxemburk J, Hynek K, Čejka T. Encrypted traffic classification: The QUIC case[C]//Proceedings of 2023 7th Network Traffic Measurement and Analysis Conference. Piscataway: IEEE, 2023: 1-10.
- [11] Nie Mingjie, Zou Futai, Qin Yi, et al. QUIC-CNN: Website fingerprinting for QUIC traffic in Tor network[C]//Proceedings of 2022 IEEE 24th International Conference on High Performance Computing & Communications. Piscataway: IEEE, 2022: 663-671.
- [12] Sirinam P, Imani M, Juarez M, et al. Deep fingerprinting: Undermining website fingerprinting defenses with deep

- learning[C]//Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2018: 1928-1943.
- [13] Chen Mantun, Wang Yongjun, Xu Hongzuo, et al. Few-shot website fingerprinting attack[J]. *Computer Networks*, 2021, 198: 108298.
- [14] Panchenko A, Lanze F, Pennekamp J, et al. Website fingerprinting at internet scale[C]//Proceedings of the 23rd Annual Network and Distributed System Security Symposium. San Diego, CA, USA: The Internet Society, 2016: 1-5.
- [15] Emmanuel T, Maupong T, Mpoeleng D, et al. A survey on missing data in machine learning[J]. *Journal of Big Data*, 2021, 8(1): 236299679.
- [16] Abe K, Goto S. Fingerprinting attack on Tor anonymity using deep learning[J]. *Proceedings of the Asia-Pacific Advanced Network*, 2016, 42: 15-20.
- [17] Rimmer V, Preuveneers D, Juarez M, et al. Automated website fingerprinting through deep learning[C]//Proceedings of the 25th Annual Network and Distributed System Security Symposium. San Diego: The Internet Society, 2018: 1-15.
- [18] Wang Yanbin, Xu Haitao, Guo Zhenhao, et al. snWF: Website fingerprinting attack by ensembling the snapshot of deep learning[J]. *IEEE Transactions on Information Forensics and Security*, 2022, 17: 1214-1226.
- [19] Shen Meng, Ji Kexin, Gao Zhenbo, et al. Subverting website fingerprinting defenses with robust traffic representation[C]//Proceedings of the 32nd USENIX Security Symposium. Anaheim, CA, USA: USENIX Association, 2023: 607-624.
- [20] Deng Xinhao, Li Qi, Xu Ke. Robust and reliable early-stage website fingerprinting attacks via spatial-temporal distribution analysis[C]//Proceedings of the 2024 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2024: 1997-2011.
- [21] Shen Meng, Wu Jinhe, Ai Junyu, et al. Swallow: A transfer-robust website fingerprinting attack via consistent feature learning[C]//Proceedings of the 2025 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2025: 1574-1588.
- [22] Deng Xinhao, Yin Qilei, Liu Zhuotao, et al. Robust multi-tab website fingerprinting attacks in the wild[C]//Proceedings of 2023 IEEE Symposium on Security and Privacy. Piscataway: IEEE, 2023: 1005-1022.
- [23] Oh S E, Mathews N, Rahman M S, et al. GANDaLF: GAN for data-limited fingerprinting[J]. *Proceedings on Privacy Enhancing Technologies*, 2021, 2021(2): 305-322.
- [24] Dingledine R, Mathewson N, Syverson P. Tor: The second-generation onion router[R]. Washington: Naval Research Lab, 2004.
- [25] Zhou Guangmeng, Guo Xiongwen, Liu Zhuotao, et al. TrafficFormer: An efficient pre-trained model for traffic data[C]//Proceedings of 2025 IEEE Symposium on Security and Privacy. Piscataway: IEEE, 2025: 1844-1860.
- [26] Xie Yi, Feng Jiahao, Huang Wenju, et al. Contrastive fingerprinting: A novel website fingerprinting attack over few-shot traces[C]//Proceedings of the ACM Web Conference 2024. New York: ACM, 2024: 1203-1214.
- [27] Chen Mantun, Wang Yongjun, Zhu Xiatian. Few-shot website fingerprinting attack with meta-bias learning[J]. *Pattern Recognition*, 2022, 130: 108739.
- [28] Lyu Qiuyun, Xie Huihui, Wang Wei, et al. TFAN: A task-adaptive feature alignment network for few-shot website fingerprinting attacks on Tor[J]. *Computers & Security*, 2024, 144: 103980.
- [29] Zou Hongcheng, Su Jinshu, Wei Ziling, et al. An efficient cross-domain few-shot website fingerprinting attack with Brownian distance covariance[J]. *Computer Networks*, 2022, 219: 109461.
- [30] Zhou Qiang, Wang Liangmin, Zhu Huijuan, et al. Few-shot website fingerprinting attack with cluster adaptation[J]. *Computer Networks*, 2023, 229: 109780.
- [31] Zhou Qiang, Wang Liangmin, Zhu Huijuan, et al. Joint alignment networks for few-shot website fingerprinting attack[J]. *The Computer Journal*, 2024, 67(6): 2331-2345.
- [32] Zhao Xiyuan, Deng Xinhao, Li Qi, et al. Towards fine-grained webpage fingerprinting at scale[C]//Proceedings of the 2024 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2024: 423-436.
- [33] Shen Meng, Ye Ke, Liu Xingtong, et al. Machine learning-powered encrypted network traffic analysis: A comprehensive survey[J]. *IEEE Communications Surveys & Tutorials*, 2023, 25(1): 791-824.
- [34] Zhan Pengwei, Wang Liming, Tang Yi. Website fingerprinting on early QUIC traffic[J]. *Computer Networks*, 2021, 200: 108538.
- [35] RFC 9000 QUIC A UDP-based multiplexed and secure transport[S].
- [36] Kukreja S L, Löfberg J, Brenner M J. A least absolute shrinkage and selection operator (LASSO) for nonlinear system identification[J]. *IFAC Proceedings Volumes*, 2006, 39(1): 814-819.
- [37] Witsenhausen H, Wyner A. A conditional entropy bound for a pair of discrete random variables[J]. *IEEE Transactions on Information Theory*, 1975, 21(5): 493-501.
- [38] Wei Rongzhe, Yin Haoteng, Jia Junteng, et al. Under-

standing non-linearity in graph neural networks from the perspective of Bayesian inference[C]//Proceedings of the 36th International Conference on Neural Information Processing Systems. New York: Curran Associates Inc., 2022: 2466.

- [39] 顾健华, 冯建华, 许辉阳, 等. 基于有向图与卷积网络强化学习的端侧协同算力资源分配方法[J]. 电子学报, 2025, 53(6): 1771-1783.

Gu Jianhua, Feng Jianhua, Xu Huiyang, et al. Directed graph and convolutional network reinforcement learning for terminal-side collaborative computing resource allocation scheme[J]. Acta Electronica Sinica, 2025, 53(6): 1771-1783. (in Chinese)

- [40] Cao Shaosheng, Lu Wei, Xu Qiongkai. GraRep: Learning graph representations with global structural information[C]//Proceedings of the 24th ACM International Conference on Information and Knowledge Management. New York: ACM, 2015: 891-900.
- [41] Battaglia P W, Hamrick J B, Bapst V, et al. Relational inductive biases, deep learning, and graph networks[PP/

OL]. V3. arXiv (2018-10-17)[2026-03-10]. <https://arXiv.org/abs/1806.01261>.

- [42] Gahtan B, Shahla R J, Cohen R, et al. Exploring QUIC dynamics: A large-scale dataset for encrypted traffic analysis[C]//Proceedings of 2025 IEEE International Mediterranean Conference on Communications and Networking. Piscataway: IEEE, 2025: 1-6.
- [43] Shen Meng, Ji Kexin, Wu Jinhe, et al. Real-time website fingerprinting defense via traffic cluster anonymization[C]//Proceedings of 2024 IEEE Symposium on Security and Privacy. Piscataway: IEEE, 2024: 3238-3256.
- [44] Gong Jiajun, Wang Tao. Zero-delay lightweight defenses against website fingerprinting[C]//Proceedings of the 29th USENIX Conference on Security Symposium. Berkeley: USENIX Association, 2020: 717-734.
- [45] Cai Xiang, Nithyanand R, Wang Tao, et al. A systematic approach to developing and evaluating website fingerprinting defenses[C]//Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security. New York: ACM, 2014: 227-238.

作者简介



谭静文 女, 1996年3月出生于黑龙江省哈尔滨市。现为哈尔滨工程大学计算机科学与技术学院博士研究生。主要研究方向为流量识别、网站指纹攻击与防御。
E-mail: tt19@hrbeu.edu.cn



杨武 男, 1974年10月出生于黑龙江省哈尔滨市。现为哈尔滨工程大学计算机科学与技术学院教授、博士生导师。主要研究方向为无线传感器网络、点对点网络 and 信息安全。中国电子学会会员编号: E190035076S。
E-mail: yangwu@hrbeu.edu.cn



王焕然 男, 1988年6月出生于黑龙江省哈尔滨市。现为哈尔滨工程大学计算机科学与技术学院讲师。主要研究方向为社交网络、隐私保护和表征学习。
E-mail: huanran.wang@hrbeu.edu.cn



秦克伟 男, 1983年6月出生于安徽省亳州市。现为中移系统集成有限公司华东区总经理、高级工程师。主要研究方向为微电子学和深度学习。
E-mail: qinkewei@cmict.chinamobile.com



韩帅 女, 1991年11月出生于黑龙江省哈尔滨市。现为哈尔滨工程大学计算机科学与技术学院讲师。主要研究方向为社交网络挖掘、大数据管理和安全。
E-mail: hshuai@hrbeu.edu.cn