

面向匿名代理流量检测的行为不变量状态空间模型

毛蔚轩¹, 王升², 刘玉玺³, 王鹏翩¹

(1. 国家计算机网络应急技术处理协调中心 北京 中国 100029; 2. 西安交通大学网络空间安全学院 西安 中国 710049; 3. 酒泉卫星发射中心, 甘肃 中国 735008)

摘要: 针对现有匿名代理检测方法依赖协议特定特征、跨协议泛化能力不足的问题, 提出行为不变量状态空间模型 (behavior-invariant state-space model, BI-SSM)。首先, 基于微突发序列设计上下文门控的自适应协同嵌入, 学习行为属性间的关联结构; 其次, 以行为单元的时间间隔调节状态演化步长, 并动态生成状态输入与输出投影, 刻画非均匀时间间隔下的时序依赖; 最后, 通过统计矩与分位数聚合会话状态分布, 采用分类损失与监督对比损失联合优化流级表示。实验结果表明, BI-SSM 在涵盖 4 种代理协议的 2 个数据集上开展的留一协议实验中, 平均 F1 分数 (F1-score) 达到 98.87%, 较性能最优的代理专用基线提高 2.60 个百分点, 验证了其跨协议检测能力。

关键词: 匿名代理; 流量分类; 状态空间模型; 加密流量

中图分类号: TP309

文献标志码: A

doi: XXXX

A State-Space Modeling Approach to Behavioral Invariants for Anonymous Proxy Traffic Detection

Mao Weixuan¹, Wang Sheng², Liu Yuxi³, Wang Pengpian¹

1. National Computer Network Emergency Response Technical Team/Coordination Center of China, Beijing 100029, China

2. School of Cyber Science and Engineering, Xi'an Jiaotong University, Xi'an 710049, China

3. Jiuquan Satellite Launch Center, Jiuquan, Gansu 735008, China

Abstract: Existing methods for detecting anonymous proxy traffic rely heavily on protocol-specific features and generalize poorly to unseen protocols. To improve cross-protocol generalization, a behavior-invariant state-space model (BI-SSM) is proposed. BI-SSM first uses a context-gated adaptive collaborative embedding to capture dependencies among behavioral attributes in microburst sequences. It then uses the intervals between behavioral units to control the state-evolution step size and dynamically generates state input and output projections to model temporal dependencies on a nonuniform time axis. Finally, statistical moments and quantiles are used to summarize the distribution of hidden states within each session, while classification and supervised contrastive losses jointly optimize the flow-level representation. Leave-one-protocol-out experiments on two datasets covering four proxy protocols show that BI-SSM achieves an average F1 score of 98.87%, outperforming the best proxy-specific baseline by 2.60 percentage points and demonstrating its ability to detect unseen proxy protocols.

Key words: anonymous proxy, traffic classification, state-space model, encrypted traffic

收稿日期: XXXX-XX-XX; 修回日期: XXXX-XX-XX

通信作者: 王鹏翩, wangpengpian@cert.org.cn。

基金项目: 国家自然科学基金资助项目 (No. XXXXXXXX, No. XXXXXXXX); 陕西省重点研发计划基金资助项目 (No. XXXXXXXX, No. XXXXXXXX)

Foundation Items: The National Natural Science Foundation of China (No. XXXXXXXX, No. XXXXXXXX), The Key Research and Development Program of Shaanxi Province (No. XXXXXXXX, No. XXXXXXXX)

0 引言

高级持续性威胁 (advanced persistent threat, APT) 常借助匿名代理构建隐蔽中继链路, 以转发命令与控制 (command and control, C2) 会话并隐藏真实通信端点^[1-2]。攻击者可在不同代理工具、加密封装和协议变体之间切换, 也可通过改变握手格式、填充策略和传输节奏形成新的流量形态^[3]。随着 Shadowsocks、Trojan 和 VMess 等代理工具广泛应用, 匿名代理呈现出部署门槛低、协议演进快和检测滞后明显等特点, 使依赖固定字段或协议指纹的边界检测方法难以及时适应。因此, 研究不依赖特定协议形态的匿名代理检测方法, 对发现隐蔽中继通信具有重要意义。

目前, 加密流量检测主要采用深度学习方法, 从原始字节、报文序列或流级数据中自动学习特征, 并在已知类别检测中取得了较好效果^[4-6]。然而, 这类模型以训练样本的可分性为目标, 容易将握手格式、封装开销和局部序列模式等协议相关信息作为判别依据。这些特征能够区分训练中出现过的协议, 却不能稳定反映匿名代理共有的中继行为, 因而在协议形态发生变化时容易出现性能下降。

匿名代理检测旨在从加密流量中识别代理中继行为。现有方法主要利用报文时空特征、突发统计、会话关联或流量路径进行检测, 少量研究也开始关注不同代理协议之间的分布偏移^[7-9]。但相关特征多依据已知协议中的判别效果选取, 尚未系统建立观测特征与代理中继机制之间的联系。现有研究未能充分考察对训练阶段未出现的新代理协议的检测能力, 也缺少成本受限自适应攻击验证。因此, 仍需从中继机制出发, 构建能够反映跨协议共同行为的流量表示, 并验证其对未知代理协议的检测能力。

本文认为, 不同代理协议虽然采用不同的加密算法、握手过程和封装格式, 但都需要在通信双方之间完成流量承接、调度和转发。该过程受到业务载荷、通信次序、链路传输和处理时间的共同约束, 会在交互方向、传输规模和时间调度之间形成联合统计关系。相较于非代理的直接通信, 额外的中继转发和调度过程使这些关系呈现不同的分布。此类统计时空约束由代理中继转发机制产生, 在不同代理协议间具有相对稳定性。攻击者若要显著破

坏这些约束, 通常需要付出较高的带宽或时延代价, 本文将其称为行为不变量。这里的不变性并非指单一观测量无法改变, 而是指联合关系在有限攻击预算下仍保持相对稳定。

基于上述的认识, 本文提出了行为不变量状态空间模型 (behavior-invariant state-space model, BI-SSM)。首先, 依据方向转换和时间间隔将双向流组织为微突发序列, 通过上下文门控的自适应协同嵌入学习行为属性间的关联; 其次, 以行为单元的实际时间间隔调节状态演化步长, 并根据当前行为表示动态生成状态输入与输出投影, 刻画非均匀时间轴上的中继依赖; 最后, 通过统计矩和分位数聚合会话状态分布, 结合分类损失与监督对比损失优化流级表示。与现有工作相比, 本文的主要贡献如下:

1) 界定了代理中继过程中的行为不变量, 并从跨协议稳定性和对抗可行性两方面开展量化分析。区分了底层固有约束与可篡改表层特征, 揭示了修改表层特征与破坏行为不变量所需的带宽和时延代价差异。

2) 提出行为不变量状态空间模型。通过上下文门控的自适应协同嵌入学习微突发行为属性间的关联, 并将实际时间间隔引入选择性状态演化, 刻画非均匀时间轴上的中继依赖。在此基础上, 利用状态分布聚合形成会话级表示, 结合分类损失与监督对比损失进行联合优化, 提高对未知代理协议的检测能力。

3) 在 AnonProxy2023^[8]和 LFETT2021^[10]数据集上开展留一协议和成本受限自适应攻击实验。结果表明, BI-SSM 的平均准确率和 F1 分数 (F1-score) 分别达到 98.98% 和 98.87%, 较最优代理专用基线分别提高 2.34 和 2.60 个百分点, 并分析了自适应攻击下检测性能与带宽、时延开销的关系。

1 相关工作

1.1 加密与隧道流量分类

加密协议隐藏了报文载荷及应用层语义, 但报文长度、传输方向、到达间隔和突发结构等侧信道仍具有可观测性。早期虚拟专用网络 (virtual private network, VPN) 和洋葱路由 (The Onion Router, Tor) 流量分类主要通过人工构造时空统计特征区分隧道类型或隧道内应用, 付钰等^[4]对相

关数据处理、特征构建和分类方法进行了系统总结。为减少对人工特征的依赖, ET-BERT^[5]和 YaTC^[6]分别利用上下文预训练与多层级掩码自编码学习报文表示;面向长流量序列, NetMamba^[11]将选择性状态空间模型引入流量分类, Deng 等^[12]进一步结合多模态时空表示、状态空间模型和流间对比学习提取流级特征。随着研究对象由单条序列扩展到多维关系, 刘涛涛等^[13]通过并行流量图融合报文头部与有效载荷信息, MH-Net^[14]利用多视图异构图建模不同粒度的字节关联。近期研究进一步关注动态环境中的未知流量, GCLC^[15]和 MMF^[16]分别从开放世界分类和少样本多标签网站指纹角度提升模型对新类别的适应能力。上述方法不断增强流量表示能力, 但其判别信息可能混合应用组成、协议实现和采集环境等因素, 对未见协议的泛化能力仍显不足。

1.2 匿名代理流量检测

匿名代理流量检测已由针对特定协议的统计判别逐步发展到流量序列、中继结构和跨协议表征。Wu 等^[17]从内容随机性和连接特征揭示了全加密代理协议的可检测性; He 和 Li^[7]将流前部报文的长度、方向及到达间隔编码为紧凑序列, 以轻量卷积网络分别检测 ShadowsocksR 和 VPN 流量; 武雯欣等^[18]进一步融合流量图像与时序特征, 识别采用不同伪装策略的 V2Ray 类流量。为减少对具体协议形态的依赖, Xue 等^[3,19]利用代理形成的嵌套协议栈, 从加密流中识别被封装的传输层安全 (transport layer security, TLS) 握手, 随后又以传输层与应用层往返时延 (round-trip time, RTT) 差异刻画代理引入的路径分段; ProxyKiller^[8]则将多条关联会话构造成行为图, 从连接结构中识别代理活动。面向跨协议分类, SPTC^[9]通过校准协议封装造成的长度偏移, 并利用累积和路径的高阶签名描述流量全局形态, 验证了长度表示在不同代理协议间的迁移能力。上述方法多围绕特定侧信号建立表示, 对请求转发与响应这一代理中继本质行为刻画不足, 跨协议检测能力仍有提升空间。

1.3 跨协议泛化与对抗鲁棒性

未知代理协议检测不仅取决于模型的表示能力, 还涉及协议变化下的分布偏移及评价结果的可靠性。MEMENTO^[20]采用类别增量学习适应持续出现的新应用, RoNeTC^[21]和 GCLC^[15]则分别通过

不确定性估计和图对比聚类识别开放环境中的未知流量。这些方法处理的是新增类别学习或未知类别识别, 而未知代理协议检测要求在代理标签不变的条件下跨越正类内部的协议域偏移。除协议变化外, Zhao 等^[22]发现数据处理过程可能引入与标签偶然相关的分类捷径, Wickramasinghe 等^[23]进一步指出会话字段、采集上下文和数据泄漏可能造成流量分类性能高估, 因此跨协议评价还需关注端点、采集批次和数据来源可能造成的评价偏差。攻击者也可能主动改变流量分布, Ding 等^[24]研究了面向流量分类器的定向通用攻击, NetMasquerade^[25]进一步在硬标签黑盒条件下通过修改报文序列模仿正常流量; Beyond RTT^[26]验证了调度、填充和报文重塑对住宅代理检测的影响, Securitas^[27]则通过学习引导的报文分片和伪报文插入抑制流量侧信道, 并量化了防护效果与带宽开销之间的权衡。

综上所述, 现有研究已利用多种侧信号开展匿名代理流量检测, 但尚未充分刻画请求经代理转发后返回响应这一行为过程。如何将这一过程中跨协议稳定的行为关系建模为行为不变量, 实现未知代理检测, 并在成本受限攻击下验证其稳定性, 是当前需要解决的关键问题。

2 方法

2.1 问题定义与总体框架

匿名代理检测的目标是在不解析加密载荷和协议字段的条件下, 根据边界侧观测到的双向会话判断其是否具有代理中继行为。将一条按到达时间排序的双向会话表示为:

$$\mathcal{F} = \left\{ (d_i, t_i, l_i) \right\}_{i=1}^N, \quad t_{i+1} \geq t_i \quad (1)$$

其中, N 为会话中的有效载荷报文数, $d_i \in \{-1, +1\}$ 表示第 i 个报文相对于会话发起端的传输方向, t_i 表示报文到达时间, l_i 表示传输层有效载荷长度。检测标签记为 $y \in \{0, 1\}$, 其中, $y = 1$ 表示匿名代理流量, $y = 0$ 表示非代理流量。模型不使用 IP 地址、端口、载荷字节、服务器名称指示、证书及协议字段等可能与特定代理实现或采集环境相关的信息。

本文关注代理正类内部存在协议分布偏移的二分类问题。设训练阶段可见的代理协议集合为 $\mathcal{P}_t$, 待检测代理协议为 p_u , 且 $p_u \notin \mathcal{P}_t$ 。模型仅使用 $\mathcal{P}_t$

中的代理流量和非代理流量进行训练，并在不利用 p_u 样本更新模型参数的条件下，对包含协议 p_u 的会话进行代理与非代理判定。该设置要求模型减少对已知协议封装形式的依赖，学习不同代理协议之间相对稳定的中继行为关系。

为此，提出行为不变量状态空间模型 (BI-SSM)，总体结构如图 1 所示。首先，根据报文方向变化和相邻报文到达间隔，将双向会话划分

为若干微突发，并从每个微突发中提取方向、载荷量、报文数、突发间隔和持续时间，形成微突发行序列。其次，采用自适应协同嵌入分别编码方向、传输规模和时间信息，并根据三类信息的联合上下文调整特征通道的贡献。然后，将微突发的实际时间间隔引入选择性状态空间模型^[28]的离散步长，使隐状态按照非均匀物理时间演化，并通过输入相关的状态投影刻画不同交互阶段对会话状态的影响。最后，利用统计矩和分位数聚合状态序列，得到流级行为表示；分类头输出代理检测概率，度量头通过监督对比约束提高同类表示的紧凑性和异类表示的可分性。

2.2 行为不变量

不同匿名代理协议在握手格式、加密封装和流量混淆方式上存在差异，但均需在通信双方之间完成数据承接、调度转发和响应回传。协议实现的变化可以改变单个报文的长度、发送时刻及分段方式，却不能消除业务数据在双向链路上的传输与交互过程。该过程受到业务载荷、请求一响应次序、链路传输能力和系统处理时间等因素的共同约束，并在流量中表现为方向、传输规模、分段粒度和时间过程之间的联合关系。本文将这种由中继机制产

生、在不同协议实现间相对稳定，且在有限带宽和时延开销下难以被同时改变的联合时空关系称为行为不变量。这里的“不变”并非指某个观测量保持固定，而是指多维观测量之间的关系相较于协议格式具有更强的稳定性。

由式(1)可知，单条双向会话能够提供方向、到达时间和载荷长度3类协议无关观测量。若直接以单个报文作为建模单元，表示结果容易受到传输控制协议 (Transmission Control Protocol, TCP) 分段、随机填充和局部网络抖动的影响。为此，按照方向变化和时间连续性将报文序列划分为微突发。定义边界指示变量为：

$$r_i = \begin{cases} 1, & i = 1, \\ 1, & 2 \leq i \leq N, d_i \neq d_{i-1} \vee t_i - t_{i-1} > \tau, \\ 0, & \text{其他} \end{cases} \quad (2)$$

其中， $r_i = 1$ 表示第 i 个报文是一个新微突发的起点， $r_i = 0$ 表示该报文与前一个报文属于同一微突发， τ 为时间异步阈值。所有满足 $r_i = 1$ 的位置构成微突发起始位置集合 $\mathcal{S} = \{i | r_i = 1, 1 \leq i \leq N\}$ 。将集合 $\mathcal{S}$ 中的元素按升序记为 $s_1, s_2, \dots, s_K$ ，并令 $e_k = s_{k+1} - 1$ 和 $e_k = N$ ，则会话可表示为微突发序列：

$$\mathcal{B} = \{b_k\}_{k=1}^K, \quad b_k = \{i | s_k \leq i \leq e_k\} \quad (3)$$

其中， K 为微突发数。每个微突发由方向一致且相邻到达间隔不超过 τ 的连续报文组成。方向变化通常对应双向通信阶段的转换，较长的报文间隔则用于分离时间上相对独立的传输阶段。因此，微突发用于近似描述会话中的连续传输行为，而不等同于对应用层请求一响应语义的严格恢复。

从交互、规模和时间3个机制维度出发，对第 k 个微突发提取如下特征：

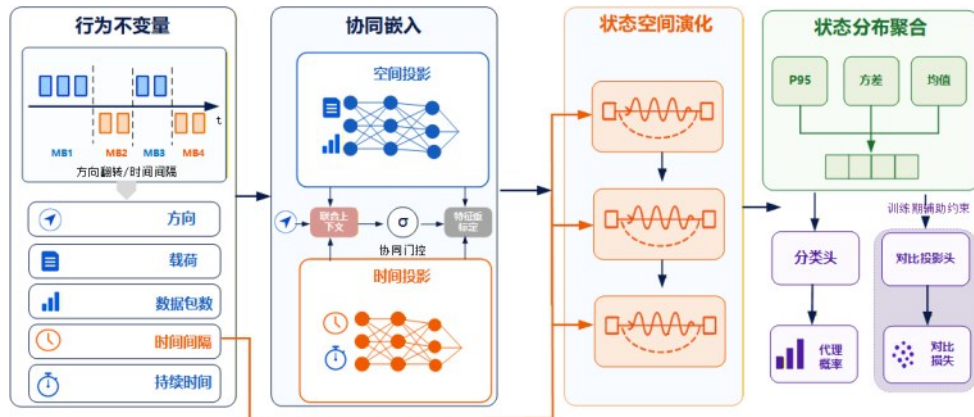


图1 行为不变量驱动的跨协议代理检测模型架构

$$\begin{aligned}
d_k &= d_{s_k}, \\
V_k &= \sum_{i=s_k}^{e_k} l_i, \\
n_k &= e_k - s_k + 1, \\
g_k &= \begin{cases} 0, & k = 1 \\ t_{s_k} - t_{e_k-1}, & k > 1 \end{cases}, \\
u_k &= t_{e_k} - t_{s_k}
\end{aligned} \quad (4)$$

其中, d_k 为微突发方向, 描述当前传输阶段相对于会话发起端的交互方向; V_k 为载荷量, 描述该阶段承载的数据规模; n_k 为报文数, 反映载荷的分段粒度; g_k 为突发间隔, 描述相邻传输阶段之间的等待时间; u_k 为持续时间, 描述当前传输阶段在时间轴上的展开过程。由此, 第 k 个微突发表示为:

$$\mathbf{x}_k = [d_k, V_k, n_k, g_k, u_k]^T \quad (5)$$

上述五维表征分别反映代理中继过程中的双向交互、传输规模、分段粒度、阶段间等待和突发内传输过程。单个观测量均可被攻击者调整, 但不同调整会引入相应的资源开销, 并可能同时改变其他维度。例如, 随机填充会直接增大载荷量, 并可能在触发载荷重分段或发送调度变化时进一步影响报文数和持续时间; 载荷重分段虽然能够改变报文数, 但会增加报文头部与协议栈处理开销; 插入反向伪报文可以改变微突发的方向转换结构, 但会增加传输字节数和双向调度开销; 主动调整突发间隔和持续时间则会增加会话完成时延。因而, 本文关注的行为不变量不是五个特征各自的取值, 而是它们在微突发序列中形成的联合关系。

在加密流量网络边界观测的既定范围内, 该表征构成了紧凑的机制行为基, 是代理中继机制所诱导的跨协议联合时空约束在微突发粒度上的可观测映射, 但不是原始流量的无损充分统计, 也不表示已经覆盖代理通信的全部行为信息。其有效性与互补性将在后续通过特征组合、方向机制消融和阈值敏感性实验进行验证。对于包含 K 个微突发的会话, 最终得到行为序列:

$$\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_K]^T \in \mathbb{R}^{K \times 5} \quad (6)$$

该序列保留微突发之间的排列和时间关系, 并作为自适应协同嵌入模块的输入。

2.3 自适应协同嵌入

行为序列同时包含离散方向变量和连续统计变量, 不同分量在数据类型、量纲和取值范围上存在

明显差异。其中, 载荷量、报文数及时间变量通常呈长尾分布, 若直接拼接输入模型, 数值较大的分量可能对表示学习产生过强影响。为统一不同分量的数值尺度, 首先对连续特征进行对数压缩和标准化处理。令 $\mathbf{q}_k = [V_k, n_k, g_k, u_k]^T$, 则第 j 个连续分量的处理结果为:

$$q_{k,j} = \frac{\log(1 + q_{k,j}) - \mu_j}{\sigma_j + \varepsilon}, \quad j \in \{1, 2, 3, 4\} \quad (7)$$

式 (7) 中, 各物理量在处理前采用统一的计量单位, μ_j 和 σ_j 分别为第 j 个对数变换特征在训练集上的均值和标准差, ε 为防止分母为 0 的常数。训练完成后, 验证集和测试集均使用训练集统计量进行变换, 避免引入测试数据的分布信息。

根据特征对应的行为属性, 将标准化后的连续特征划分为规模统计分量和时间过程分量:

$$\mathbf{x}_k^s = [q_{k,1}, q_{k,2}]^T, \quad \mathbf{x}_k^t = [q_{k,3}, q_{k,4}]^T \quad (8)$$

其中, $\mathbf{x}_k^s$ 由载荷量和报文数组成, 用于描述微突发的传输规模及分段粒度; $\mathbf{x}_k^t$ 由突发间隔和持续时间组成, 用于描述微突发之间及其内部的时间过程。方向 d_k 作为离散变量单独编码。三类分量分别映射至隐空间:

$$\mathbf{e}_k^d = \text{Emb}(d_k), \mathbf{e}_k^s = \phi_s(\mathbf{x}_k^s), \mathbf{e}_k^t = \phi_t(\mathbf{x}_k^t) \quad (9)$$

其中, $\text{Emb}(\cdot)$ 为方向嵌入函数, $\phi_s(\cdot)$ 和 $\phi_t(\cdot)$ 分别为规模统计分量与时间过程分量的可学习投影函数。将三类嵌入拼接, 得到第 k 个微突发的联合上下文表示:

$$\mathbf{c}_k = [\mathbf{e}_k^d; \mathbf{e}_k^s; \mathbf{e}_k^t] \in \mathbb{R}^{d_c} \quad (10)$$

方向、规模和时间分量在不同代理协议及通信阶段中的判别作用并不固定。例如, 相近的载荷规模可能对应不同的交互方向和等待过程, 仅依赖单一分量难以描述这种条件关联。为此, 利用联合上下文生成通道门控权重:

$$\mathbf{a}_k = \sigma(W_2 \text{GELU}(W_1 \mathbf{c}_k + \mathbf{b}_1) + \mathbf{b}_2) \quad (11)$$

其中, W_1 、 W_2 、 $\mathbf{b}_1$ 和 $\mathbf{b}_2$ 为可学习参数, $\text{GELU}(\cdot)$ 为高斯误差线性单元激活函数, $\sigma(\cdot)$ 为 Sigmoid 函数, $\mathbf{a}_k \in (0, 1)^{d_c}$ 与 $\mathbf{c}_k$ 具有相同维度。由于每个通道的权重均由三类分量的联合表示生成, 门控过程能够根据当前微突发的整体行为上下文调整不同特征通道的贡献。

经过通道重标定后, 第 k 个微突发的嵌入表

示为:

$$\mathbf{z}_k = \mathbf{W}_o(\mathbf{a}_k \odot \mathbf{c}_k) + \mathbf{b}_o \quad (12)$$

其中, $\odot$ 表示逐元素乘法, $\mathbf{W}_o$ 和 $\mathbf{b}_o$ 为输出投影参数, $\mathbf{z}_k \in \mathbb{R}^{d_z}$ 。该过程并不单独将某一维特征作为判别依据, 而是在方向、规模和时间分量的联合条件下形成微突发表示, 以保留代理行为不变量所体现的多维协同关系。

对于包含 K 个微突发的会话, 自适应协同嵌入模块输出:

$$\mathbf{Z} = [\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_K]^T \in \mathbb{R}^{K \times d_z} \quad (13)$$

该序列作为选择性状态空间演化模块的输入, 用于进一步建模微突发之间的非均匀时序依赖。

2.4 选择性状态空间演化

由于网络报文并非等间隔到达, 不同微突发之间的时间间隔反映了链路传输、队列等待和系统处理等因素形成的综合时间代价。选择性状态空间模型能够根据输入内容调整状态更新参数, 但其离散步长不一定与实际突发间隔保持一致。为此, 本文在输入选择性建模的基础上引入物理时间约束, 使状态演化同时受微突发内容和实际时间间隔控制。连续状态空间模型表示为:

$$\begin{aligned} \dot{\mathbf{h}}(t) &= \mathbf{A}\mathbf{h}(t) + \mathbf{B}(t)\mathbf{z}(t) \\ \mathbf{o}(t) &= \mathbf{C}(t)\mathbf{h}(t) + \mathbf{D}_z\mathbf{z}(t) \end{aligned} \quad (14)$$

其中, $\mathbf{z}(t)$ 为输入表示, $\mathbf{h}(t)$ 为隐状态, $\mathbf{o}(t)$ 为状态输出, $\mathbf{A}$ 为状态转移矩阵, $\mathbf{B}(t)$ 和 $\mathbf{C}(t)$ 分别为输入投影和输出投影, $\mathbf{D}_z$ 为直接传递矩阵。为保证连续状态随时间稳定演化, 将状态转移矩阵参数化为 $\mathbf{A} = -\text{diag}(\text{softplus}(\mathbf{a}_A))$, 使其对角元素保持为负值。

对于第 k 个微突发, 离散步长由当前嵌入和实际突发间隔共同确定:

$$\delta_k = \text{softplus}(\mathbf{w}_\delta^T \mathbf{z}_k + \mathbf{b}_\delta + \rho \tilde{g}_k) \quad (15)$$

其中 $\rho = \text{softplus}(a_\rho)$, $\tilde{g}_k$ 为标准化后的突发间隔, $\mathbf{w}_\delta$ 、 $\mathbf{b}_\delta$ 和 a_ρ 为可学习参数。上式保留了选择性状态空间模型根据输入内容调整步长的能力, 同时通过非负参数 ρ 显式建立实际时间间隔与状态演化尺度之间的联系。突发间隔在第 2.3 节中作为行为属性参与协同嵌入, 在本节中则直接控制状态转移尺度, 二者分别作用于输入表征和历史状态演化。

为使模型根据当前微突发选择需要写入和读取的状态信息, 输入投影和输出投影由 $\mathbf{z}_k$ 动态生成:

$$\begin{aligned} \mathbf{B}_k &= \text{reshape}(\mathbf{W}_B \mathbf{z}_k + \mathbf{b}_B) \\ \mathbf{C}_k &= \text{reshape}(\mathbf{W}_C \mathbf{z}_k + \mathbf{b}_C) \end{aligned} \quad (16)$$

其中, $\mathbf{W}_B$ 、 $\mathbf{W}_C$ 、 $\mathbf{b}_B$ 和 $\mathbf{b}_C$ 为不同时间步共享的可学习参数, 生成的 $\mathbf{B}_k$ 和 $\mathbf{C}_k$ 则随微突发内容变化。该选择性参数化使不同交互阶段以不同强度影响隐状态, 避免对所有微突发采用固定的状态写入和读取方式。

采用零阶保持方法, 将连续状态空间模型按照 δ_k 离散化为:

$$\begin{aligned} \bar{\mathbf{A}}_k &= \exp(\delta_k \mathbf{A}) \\ \bar{\mathbf{B}}_k &= \left(\int_0^{\delta_k} \exp(\eta \mathbf{A}) d\eta \right) \mathbf{B}_k \end{aligned} \quad (17)$$

并沿微突发序列执行状态更新:

$$\begin{aligned} \mathbf{h}_k &= \bar{\mathbf{A}}_k \mathbf{h}_{k-1} + \bar{\mathbf{B}}_k \mathbf{z}_k \\ \mathbf{o}_k &= \mathbf{C}_k \mathbf{h}_k + \mathbf{D}_z \mathbf{z}_k \end{aligned} \quad (18)$$

其中, 初始状态 $\mathbf{h}_0 = \mathbf{0}$ 。由于 $\mathbf{A}$ 采用负对角参数化, 较大的 δ_k 对应更强的历史状态衰减, 从而使长时间等待后的微突发与连续到达的微突发产生不同的状态更新结果。最终得到状态输出序列 $\mathbf{O} = [\mathbf{o}_1, \mathbf{o}_2, \dots, \mathbf{o}_K]^T$, 用于后续的流级状态分布聚合。

2.5 状态分布聚合与联合优化

代理中继行为通常分布在会话的多个交互阶段, 若仅使用末端状态进行判定, 结果容易受到会话长度、结束阶段及局部异常状态的影响。为获得稳定的流级表示, 本文对整个状态序列进行分布聚合。

首先, 沿微突发维度计算状态序列的一阶矩和二阶中心矩:

$$\begin{aligned} \boldsymbol{\mu}_o &= \frac{1}{K} \sum_{k=1}^K \mathbf{o}_k \\ \boldsymbol{\sigma}_o^2 &= \frac{1}{K} \sum_{k=1}^K (\mathbf{o}_k - \boldsymbol{\mu}_o) \odot (\mathbf{o}_k - \boldsymbol{\mu}_o) \end{aligned} \quad (19)$$

其中, $\boldsymbol{\mu}_o$ 描述会话状态的整体中心, $\boldsymbol{\sigma}_o^2$ 描述不同交互阶段的状态波动。考虑到代理协议中的握手、控制交互及局部突发可能产生非对称状态分布, 进一步沿微突发维度逐通道提取分位数:

$$\mathbf{q}_\alpha = \text{Quantile}_k(\{\mathbf{o}_k\}_{k=1}^K, \alpha) \quad (20)$$

其中 $\alpha \in \{0.25, 0.50, 0.75, 0.90\}$, 分位数能够补充均值和方差对偏态分布及高响应状态描述不足的问题, 同时降低个别极端状态对聚合结果的影响。将统计矩与分位数拼接并进行投影, 得到流级行为

表示:

$$\mathbf{v} = \text{LN}(\mathbf{W}_v[\boldsymbol{\mu}_o; \boldsymbol{\sigma}_o^2; \mathbf{q}_{0.25}; \mathbf{q}_{0.50}; \mathbf{q}_{0.75}; \mathbf{q}_{0.90}] + \mathbf{b}_v) \quad (21)$$

其中, $\mathbf{W}_v$ 和 $\mathbf{b}_v$ 为可学习参数, $\text{LN}(\cdot)$ 表示层归一化。对于长度不同的微突发序列, 掩码位置不参与上述统计量计算。

在流级表示基础上, 采用分类头和度量头进行联合优化。分类头输出会话属于匿名代理流量的概率:

$$\hat{p}_i = \sigma(\mathbf{w}_c^T \mathbf{v}_i + b_c) \quad (22)$$

其中, $\mathbf{w}_c$ 和 b_c 为分类参数, $\sigma(\cdot)$ 为 Sigmoid 函数。对于包含 M 个样本的训练批次, 分类损失采用二元交叉熵:

$$\mathcal{L}_{\text{ce}} = -\frac{1}{M} \sum_{i=1}^M [y_i \log \hat{p}_i + (1 - y_i) \log (1 - \hat{p}_i)] \quad (23)$$

分类损失用于学习代理与非代理之间的判定边界, 但不直接约束同类样本在表示空间中的分布。为减少同一类别内部由协议差异造成的表示离散, 度量头将流级表示投影至单位超球面:

$$\mathbf{r}_i = \frac{\psi(\mathbf{v}_i)}{\|\psi(\mathbf{v}_i)\|_2} \quad (24)$$

其中, $\psi(\cdot)$ 为可学习投影函数。训练阶段对流级表示 $\mathbf{v}_i$ 施加轻量随机扰动以形成增强视图, 该操作用于表示空间正则化。同标签样本以及同一原始会话的增强视图共同构成正样本集合, 并采用监督对比损失进行度量约束^[29]。首先定义样本间相似度及其归一化项:

$$s_{ij} = \mathbf{r}_i^T \mathbf{r}_j \\ Z_i = \sum_{a \in \mathcal{A}(i)} \exp(s_{ia}/\kappa) \quad (25)$$

则监督对比损失表示为:

$$\mathcal{L}_{\text{con}} = -\frac{1}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} \frac{1}{|\mathcal{P}(i)|} \sum_{p \in \mathcal{P}(i)} \log \frac{\exp(s_{ip}/\kappa)}{Z_i} \quad (26)$$

其中, s_{ij} 表示样本 i 与样本 j 的余弦相似度, Z_i 为锚样本 i 对应的归一化项, $\mathcal{I}$ 为锚样本集合, $\mathcal{P}(i)$ 为与样本 i 标签相同的正样本集合, $\mathcal{A}(i)$ 为除样本 i 外的对比样本集合, κ 为温度系数。不同已知代理协议的样本具有相同代理标签, 因而该损失能够显式约束其流级表示的类内紧凑性, 并保持代理与非代理表示之间的可分性。模型最终采用联合目标进行训练:

$$\mathcal{L} = \mathcal{L}_{\text{ce}} + \lambda \mathcal{L}_{\text{con}} \quad (27)$$

其中, λ 为两项损失的权衡系数。分类头负责学习代理检测边界, 度量头用于约束流级表示的空间结构。推理阶段仅使用分类头输出的 $\hat{p}_i$ 完成代理与非代理流量判定。

3 实验

3.1 数据集与实验设置

实验采用 AnonProxy2023^[8] 和 LFETT2021^[10] 两个公开数据集。AnonProxy2023 包含四种不同匿名代理协议产生的加密转发流量, LFETT2021 则选取 VMess 和 ShadowsocksR 作为补充, 共选取 Shadowsocks、ShadowsocksR、Trojan 和 VMess 会话作为代理正类, 并以正常通信会话作为非代理负类, 构建代理与非代理二分类任务。经过统一预处理后, 共得到 282979 条双向会话, 各类样本构成如表 1 所示。

表 1 实验样本构成			
类别	协议	会话数	占比/%
代理	Shadowsocks	32480	11.48
代理	ShadowsocksR	29615	10.47
代理	Trojan	27944	9.87
代理	VMess	35821	12.66
代理	合计	125860	44.48
非代理	正常通信	157119	55.52
合计	-	282979	100

在数据预处理阶段, 首先重组 IP 分片, 再按照五元组建立双向 TCP 会话。以首个同步序列号 (synchronize sequence numbers, SYN) 报文的发送端作为会话发起端; 对于缺少完整握手过程的会话, 以首个有效载荷报文的发送端确定方向。TCP 报文按照采集时间排序, 并根据序列号去除完全重传的载荷区间; 对于重叠载荷区间, 保留首次观测到的部分。使用 60s 空闲超时切分长连接, 去除纯确认报文, 仅保留传输层有效载荷长度大于 0 的报文。为排除信息不足或采集异常的样本, 过滤仅包含单向有效载荷、有效载荷报文数少于 10、持续时间小于 10ms 或大于 600s 的会话。对于位于原始数据包文件起止边界且缺少完整双向交互的截断会话, 也予以排除。数据划分前, 根据报文方向、有

效载荷长度和相对到达间隔序列计算会话摘要并删除完全重复的样本,避免相同流量序列同时进入训练集和测试集。本文实验仅评价基于 TCP 的匿名代理流量,用户数据报协议 (User Datagram Protocol, UDP) 和快速 UDP 互联网连接 (Quick UDP Internet Connections, QUIC) 流量不纳入当前实验范围。

按照第 2.2 节的方法,以方向变化或相邻报文间隔超过 $\tau = 0.5$ 为边界构造微突发。第 3.3 节进一步在 0.1s、0.2s、0.5s、1.0s 和 2.0s 条件下分析阈值敏感性,并仅根据可见协议验证集选择默认值。模型输入的最大微突发数设置为 $K_{\max} = 128$ 。当 $K > 128$ 时,保留会话前 128 个微突发,否则用掩码补齐,补齐位置不参与状态更新和分布聚合。连续特征的对数变换及标准化参数仅根据训练集计算,并直接应用于验证集和测试集。

为评价模型对未知代理协议的检测能力,采用留一协议方式开展 4 轮实验。首先,将每种代理协议和非代理样本分别按照 7:1:2 的比例随机划分为训练集、验证集和测试集。每轮选择一种代理协议作为未知协议,其训练集和验证集均不参与模型训练及参数选择。模型仅使用其余 3 种代理协议和非代理样本的训练集进行训练,并使用相应验证集确定模型参数和分类阈值。测试时,其余 3 种代理协议的测试分区用于评价已知协议检测性能,被留出协议的测试分区用于评价未知协议检测性能。每轮分别从非代理测试分区中固定抽取样本,使两类测试集中的代理样本占比均为 45%,且所有方法使用相同的抽样索引。依次留出 Shadowsocks、ShadowsocksR、Trojan 和 VMess,各项指标首先在每轮测试集上独立计算,最终报告 4 轮结果的算术平均值。

两个公开数据集均由常见代理工具和正常网络应用构建,不包含经过确认的真实 APT 代理 C2 会话。因此,本文实验直接验证的是未知匿名代理协议检测能力,以及该能力在流量扰动和成本受限规避条件下的稳定性,不能等同于对完整 APT 攻击链或真实 APT 组织流量的直接检测效果。

3.2 对比实验

为验证所提方法的检测性能,选取 ET-BERT^[5]、YaTC^[6]、NetMamba^[30]、MH-Net^[14]、He-Li^[7]和 SPTC^[9]等先进方法进行对比试验。He-Li

将流前部报文的长度、方向和到达间隔编码为定长序列,并使用一维卷积网络完成检测;SPTC 校准协议封装引起的长度偏移,通过累积和路径及其高阶签名构建流级表示。适配过程中保留两种方法的特征变换与主体模型,仅将输出调整为代理与非代理二分类,未明确的数据相关参数由可见协议验证集确定。ET-BERT 保留预训练报文编码器并将分类层替换为二分类输出;YaTC 保留多层次流表示及预训练编码结构,在本文数据上微调分类头;NetMamba+沿用原有报文序列预处理和状态空间骨干网络,将输出层调整为二分类;MH-Net 按照原方法构建多视图异构表示,并将图分类结果映射为代理检测概率。所有方法采用相同的数据划分和留一协议设置。未知协议样本不参与模型训练、参数选择和分类阈值确定。

考虑到已有研究指出^[23],流量分类模型可能利用服务器名称指示 (server name indication, SNI)、IP 首部校验和、TCP 序号及时间戳高位等非语义字段形成过拟合捷径,进而造成异常偏高的评价结果,本文参照相关遮蔽策略^[12],对强标识信息、SNI 及会话相关上下文字段进行清洗或随机化处理,避免模型依赖与代理行为无关的泄露特征。

采用准确率 (accuracy, Acc)、精确率 (Precision)、召回率 (Recall)、F1 值和受试者工作特征曲线下面积 (area under the receiver operating characteristic curve, ROC-AUC) 作为评价指标,其中代理流量为正类。Precision 反映检测告警的可靠性,Recall 反映对代理流量的检出能力,F1 综合衡量 Precision 与 Recall,ROC-AUC 用于评价不同分类阈值下模型的整体判别能力。各项结果均为 4 轮留一协议实验的算术平均值。

由表 2 可知,各方法在已知协议上均具有较好的检测能力,其中 BI-SSM 的 Acc、Precision、Recall、F1 和 ROC-AUC 分别达到 99.97%、99.94%、99.98%、99.96% 和 99.99%。与性能最好的对比方法 SPTC 相比,BI-SSM 的 Acc 和 F1 分别提高 1.00 和 1.11 个百分点。其 Precision 和 Recall 均接近 100%,表明模型能够兼顾代理流量检出能力与告警可靠性。

MH-Net 在已知协议上的 Acc 和 F1 分别达到 98.33% 和 98.15%,说明多视图字节关系能够有效

表2 已知协议检测结果(%)

方法	Acc	Precision	Recall	F1	ROC-AUC
He-Li	98.25	97.81	98.30	98.05	99.12
SPTC	98.97	98.70	99.00	98.85	99.56
ET-BERT	91.28	89.82	90.93	90.37	96.14
NetMamba+	95.56	94.73	95.45	95.09	98.07
MH-Net	98.33	97.91	98.39	98.15	99.48
YaTC	93.80	92.66	93.65	93.15	97.32
BI-SSM	99.97	99.94	99.98	99.96	99.99

描述训练阶段已经出现的协议特征。He-Li 和 SPTC 分别通过前部报文时空特征和长度路径签名获得较高检测性能,表明代理流量的序列形态对于已知协议识别具有较强判别能力。

由表 3 可知,当测试协议在训练阶段不可见时,各方法的检测性能均有所下降。BI-SSM 的 Acc、Precision、Recall、F1 和 ROC-AUC 分别达到 98.98%、98.62%、99.12%、98.87% 和 99.42%,均优于对比方法。较高的 Precision 和 Recall 表明该方法能够兼顾告警可靠性与未知代理流量检出能力,ROC-AUC 则说明模型在不同分类阈值下仍具有较强的整体判别能力。

表3 未知协议检测结果(%)

方法	Acc	Precision	Recall	F1	ROC-AUC
He-Li	94.33	93.31	94.10	93.70	97.15
SPTC	96.64	96.10	96.45	96.27	98.21
ET-BERT	86.12	84.20	85.13	84.66	90.42
NetMamba+	92.71	91.75	92.07	91.91	95.64
MH-Net	84.76	82.70	83.63	83.16	89.55
YaTC	86.94	85.10	86.04	85.57	91.38
BI-SSM	98.98	98.62	99.12	98.87	99.42

SPTC 是性能最好的对比方法,其上述 5 项指标分别为 96.64%、96.10%、96.45%、96.27% 和 98.21%;BI-SSM 相应提高 2.34、2.52、2.67、2.60 和 1.21 个百分点。进一步比较表 2 和表 3 可知,6 种对比方法从已知协议到未知协议的 F1 降幅为 2.58~14.99 个百分点,而 BI-SSM 仅下降 1.09 个百分点。结果表明,联合描述交互方向、传输规模、分段粒度及突发内外时间关系,并对其状态演化过

程进行建模,有助于减弱协议变化造成的分布偏移,提高模型对未知代理协议的泛化能力。

3.3 行为不变量表征分析

为分析多维行为表征的跨协议属性及三类行为信息的互补作用,从跨协议分布、特征组合和微突发阈值 3 个方面进行分析。

首先,在固定测试集中对各类别进行等量抽样,将逐包长度一方向一到达间隔表征与本文五维微突发表征分别采用均值、标准差和分位数汇聚为流级统计向量。所有特征均采用训练集统计量进行标准化,并使用径向基函数核最大均值差异平方(squared maximum mean discrepancy, MMD²)度量分布距离,核带宽由中值启发式确定^[31]。定义代理协议间差异与代理—非代理类别间差异之比为:

$$R_{\text{inv}} = \frac{\frac{2}{M(M-1)} \sum_{p < q} \text{MMD}^2(\mathcal{P}_p, \mathcal{P}_q)}{\frac{1}{M} \sum_{p=1}^M \text{MMD}^2(\mathcal{P}_p, \mathcal{P}_N)} \quad (28)$$

其中, $M = 4$, $\mathcal{P}_p$ 和 $\mathcal{P}_N$ 分别表示第 p 种代理协议与非代理流量的特征分布。 R_{inv} 越小,表示不同代理协议之间的分布差异相对于代理—非代理类别差异越小。

经过 10 次独立重采样,逐包基础表征的代理协议间 MMD² 和代理—非代理 MMD² 分别为 0.176 和 0.311,按式 (28) 计算得到比值为 0.566;五维微突发表征的对应结果分别为 0.083 和 0.397,比值为 0.209。相较逐包基础表征,五维微突发表征在缩小不同代理协议分布差异的同时扩大了代理与非代理流量的分布距离,说明其能够抑制部分协议实现差异,并保留代理中继行为的类别判别信息。

进一步通过特征组合和方向机制消融实验考察三类行为信息的作用,结果如表 4 所示。其中, D 、 S 和 T 分别表示方向、规模统计(V, n)和时间过程(g, u), $\Delta F1$ 表示相对于完整五维表征的 F1 变化。

由表 4 可知,单独使用方向、规模或时间特征时, F1 分别为 84.31%、90.62% 和 87.54%,均明显低于完整表征,说明单一类型的行为信息不足以稳定识别未知代理协议。任意两组特征联合后检测性能均有所提高,其中 $D + S$ 取得 95.67% 的 F1 和 97.86% 的 ROC-AUC,但与完整表征相比仍分别下降 3.20 和 1.56 个百分点。这表明交互方向、传输规

表4 行为特征组合及方向机制消融结果(%)

特征配置	F1	ROC-AUC	$\Delta F1$
D	84.31	89.74	-14.56
S	90.62	94.36	-8.25
T	87.54	92.31	-11.33
D+S	95.67	97.86	-3.20
D+T	93.61	96.18	-5.26
S+T (无 d_k)	95.21	97.42	-3.66
S+T (无方向边界)	92.76	95.74	-6.11
完整表征	98.87	99.42	0

模和时间过程之间具有明显的互补关系, 五维表征的有效性来源于多类行为信息的联合约束, 而非某个单独统计量。

当保留方向参与的微突发划分、但不将 d_k 输入模型时, $S+T$ 组合的F1为95.21%, 较完整表征下降3.66个百分点。进一步取消方向变化这一分段条件后, F1降至92.76%, ROC-AUC降至95.74%。相较仅移除 d_k , 方向无关分段使F1进一步下降2.45个百分点, 说明交互方向不仅通过显式特征参与建模, 还通过微突发边界反映请求—转发—响应过程的结构变化。完整五维表征取得98.87%的F1和99.42%的ROC-AUC, 验证了方向、规模、分段粒度及突发内外时间关系联合建模的必要性。

最后, 设置 $\tau \in \{0.1, 0.2, 0.5, 1.0, 2.0\}$, 考察微突发划分阈值对表征粒度和检测性能的影响, 结果如图2和图3所示。当 τ 由0.1s增大至2.0s时, 每条流的平均微突发数由43.8下降至16.2。在0.2~1.0s范围内, F1和ROC-AUC的最大波动分别为0.22和0.13个百分点, 说明模型对中间范围内的阈值变化不敏感。阈值过小时, 相邻传输过程容易被过度切分; 阈值过大时, 多个独立交互可能被合并。本文取 $\tau = 0.5s$, 此时F1和ROC-AUC分别达到98.87%和99.42%。上述结果说明五维表征在本文观测范围内具有较好的有效性和紧凑性。

3.4 模型消融实验

为验证各模型模块的作用, 在保持五维行为表征、数据划分和训练参数不变的条件下进行消融实验。去除自适应协同嵌入时, 采用线性投影替代门控融合; 去除选择性状态演化时, 将嵌入序列直接送入分布聚合层; 去除状态分布聚合时, 以最后一个有效状态作为流级表示; 去除对比损失时, 令度

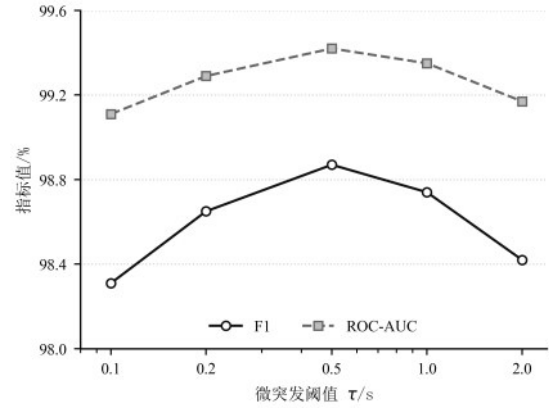


图2 不同微突发阈值下的未知协议检测性能

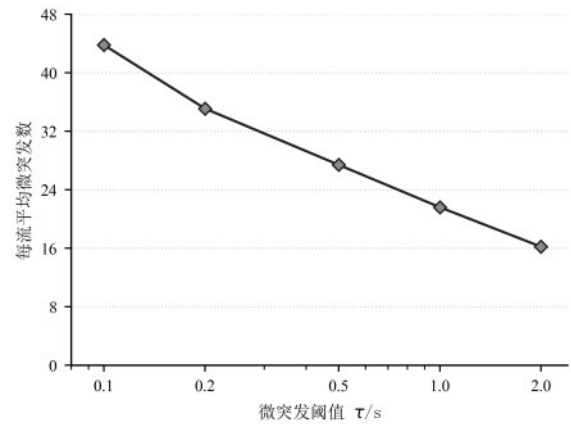


图3 不同微突发阈值下的每流平均微突发数

量损失权重 $\lambda = 0$, 仅采用分类损失训练。结果为4轮留一协议实验的算术平均值。

由表5可知, 移除各模块均会导致检测性能下降, 且未知协议上的下降更为明显。去除自适应协同嵌入后, 未知协议Acc和F1分别下降1.64和1.83个百分点, 降幅最大, 说明不同类型行为特征之间的条件化融合有助于缓解协议实现差异造成的表征偏移。去除选择性状态演化和状态分布聚合后, 未知协议F1分别下降1.52和1.19个百分点, 表明代理中继行为需要同时刻画非均匀时间轴上的状态迁移及完整序列的状态分布。

去除对比损失后, 未知协议F1下降0.84个百分点, 说明度量约束能够通过增强类内聚合与类间分离进一步改善未知协议判别。4种消融在已知协议上的F1降幅为0.26~0.61个百分点, 而在未知协议上的降幅为0.84~1.83个百分点, 表明各模块对跨协议条件下的表示稳定性具有更为明显的作用。

表5 模型消融实验结果(%)

模型配置	已知 Acc	已知 F1	未知 Acc	未知 F1
去除自适应协同嵌入	99.42	99.35	97.34	97.04
去除选择性状态演化	99.52	99.47	97.62	97.35
去除状态分布聚合	99.62	99.58	97.91	97.68
去除对比损失	99.73	99.70	98.23	98.03
完整模型	99.97	99.96	98.98	98.87

3.5 扰动与成本受限实验

为评价模型对流量表层变化的稳定性,首先沿用原稿设置,在模型参数和分类阈值不变的条件下,分别对测试集中的代理流量施加随机填充、时序扰动和突发插入。随机填充改变部分报文的载荷长度;时序扰动对选定报文施加非负延迟,同时保持报文次序不变;突发插入在原序列中加入不承载应用语义的低载荷交互。扰动强度表示被随机选中并实施相应操作的报文或微突发占比,不等于带宽或时延开销。实验结果如表6所示。

由表6可知,随着扰动强度增加,BI-SSM的F1呈平稳下降趋势。在20%扰动强度下,随机填充、时序扰动和突发插入分别使F1下降0.75、1.14和1.51个百分点。其中,随机填充的影响最小,说明模型未单纯依赖载荷长度;时序扰动和突发插入会改变微突发序列及其状态演化轨迹,因此影响相对明显。总体来看,随机作用于局部报文或微突发的变化尚不足以破坏完整流中的联合行为关系,状态分布聚合也能够减弱局部异常片段对流级判定的影响。

表6 不同随机扰动下的未知协议F1(%)

扰动强度/%	随机填充	时序扰动	突发插入
0	98.87	98.87	98.87
5	98.76	98.69	98.60
10	98.58	98.42	98.23
15	98.37	98.09	97.82
20	98.12	97.73	97.36

上述实验未考虑攻击者根据模型输出主动选择操作位置和扰动方式。为进一步评价模型面对针对性规避时的性能,构造成本受限的自适应攻击^[24-27]。假设攻击者了解模型的特征构造和参数,并能够获得输出置信度;攻击者可以实施载荷填

充、非负延迟、有效载荷重分段和伪突发插入,但不能删除或重排真实有效载荷,也不能改变真实报文的传输方向,以保证TCP会话有效和应用语义不变。

设原始流和扰动后流的总传输字节数分别为 B 和 B' ,持续时间分别为 T 和 T' ,定义带宽开销 $C_B = (B' - B)/B$ 和时延开销 $C_T = (T' - T)/T$ 。攻击目标为:

$$a^* = \arg \max_{a \in \Omega(\beta, \gamma)} \mathcal{L}_{cc}(f_{\theta}(\mathcal{T}_a(F)), 1) \quad (29)$$

其中, $C_B \leq \beta$, $C_T \leq \gamma$, $\mathcal{T}_a$ 表示对流量实施操作 a , $\Omega(\beta, \gamma)$ 表示满足带宽预算 β 和时延预算 γ 的可行操作集合。攻击过程逐步生成可行候选操作,并选择单位成本下代理置信度下降最大的操作,直至达到预算上限或分类结果发生改变,每条流最多进行100次候选评估。

攻击仅作用于代理测试样本,非代理样本和验证集确定的分类阈值均保持不变。除F1外,采用攻击成功率(attack success rate, ASR)评价规避效果,将其定义为干净测试中被正确检测、经扰动后被误判为非代理的样本比例。选取SPTC和NetMamba+进行对照,各模型分别依据自身输出生成攻击序列。实验统一令 $\beta = \gamma \in \{0\%, 5\%, 10\%, 20\%\}$ 。结果如表7所示,其中实际带宽开销和实际时延开销为BI-SSM在代理测试样本上的平均值,ASR越低表示模型越不易被规避。

由表7可知,自适应攻击能够根据模型反馈集中修改对判定影响较大的报文和微突发,其攻击效果明显强于随机扰动。随着预算上限由5%提高至20%,BI-SSM的F1由97.21%下降至91.76%,ASR由3.27%上升至13.29%,说明模型面对主动规避时仍存在性能退化,不能视为对自适应攻击完全免疫。

在20%预算下,SPTC和NetMamba+的F1分别降至84.76%和77.82%,ASR分别达到20.76%和25.09%;BI-SSM的F1分别高出7.00和13.94个百分点。相较无攻击条件,3种模型的F1分别下降11.51、14.09和7.11个百分点。该结果表明,联合利用方向、规模、分段粒度及时间过程进行判别,能够降低单一类型扰动对检测结果的影响。攻击者需要同时改变多个行为维度才能取得更高规避效

果，并由此承担相应的带宽和时延成本。

4 结束语

针对匿名代理协议变化导致表层特征失效和未知协议检测性能下降的问题，本文提出行为不变量状态空间模型 BI-SSM。该模型将双向会话组织为微突发序列，以方向、载荷体积、报文数量及突发内外时间描述代理中继机制诱导的联合时空约束，并通过协同嵌入、选择性状态演化和状态分布聚合形成流级表示。留一协议实验表明，模型在未知协议上的 Acc、Precision、Recall、F1 和 ROC-AUC 分别为 98.98%、98.62%、99.12%、98.87% 和 99.42%，F1 较最优对比方法提高 2.60 个百分点。特征与模块消融验证了各组成部分的有效性；在 20% 成本预算下，模型面对自适应规避仍取得 91.76% 的 F1。当前实验仅覆盖 4 种 TCP 代理协议，且不包含真实 APT 代理 C2 会话。未来将结合真实网络和受控攻防流量，研究模型在更多传输协议、采集环境及规避策略下的适用性。

表 7

成本受限自适应攻击结果(%)

预算上限 $\beta = \gamma$	BI-SSM 实际带宽开销	BI-SSM 实际时延开销	SPTC F1/ASR	NetMamba+ F1/ASR	BI-SSM F1/ASR
0%	0	0	96.27/0	91.91/0	98.87/0
5%	4.63	4.21	93.62/5.18	88.43/6.79	97.21/3.27
10%	9.21	8.47	90.18/11.53	84.17/14.54	95.08/7.31
20%	18.54	17.02	84.76/20.76	77.82/25.09	91.76/13.29

参考文献:

- [1] HUTCHINS E M, CLOPPERT M J, AMIN R M. Intelligence-driven computer network defense informed by analysis of adversary campaigns and intrusion kill chains[C]//Proceedings of the 6th International Conference on Information Warfare and Security. Washington, DC, USA, 2011-03-17—2011-03-18. Reading: Academic Conferences and Publishing International, 2011: 113-125.
- [2] ALSHAMRANI A, MYNENI S, CHOWDHARY A, et al. A survey on advanced persistent threats: techniques, solutions, challenges, and research opportunities[J]. IEEE Communications Surveys & Tutorials, 2019, 21(2): 1851-1877.
- [3] XUE D, KALLITSIS M, HOUMANSADR A, et al. Fingerprinting obfuscated proxy traffic with encapsulated TLS handshakes[C]//Proceedings of the 33rd USENIX Security Symposium. Philadelphia, PA, USA, 2024-08-14—2024-08-16. Berkeley: USENIX Association, 2024: 2689-2706.
- [4] 付钰, 刘涛涛, 王坤, 等. 基于机器学习的加密流量分类研究综述[J]. 通信学报, 2025, 46(1): 167-191.
- [5] LIN X J, XIONG G, GOU G P, et al. ET-BERT: a contextualized datagram representation with pre-training transformers for encrypted traffic classification[C]//Proceedings of the ACM Web Conference 2022. Lyon, France, 2022-04-25—2022-04-29. New York: ACM Press, 2022: 633-642.
- [6] ZHAO R J, ZHAN M W, DENG X W, et al. Yet another traffic classifier: a masked autoencoder based traffic transformer with multi-level flow representation[C]//Proceedings of the 37th AAAI Conference on Artificial Intelligence. Washington, DC, USA, 2023-02-07—2023-02-14. Palo Alto: AAAI Press, 2023, 37(4): 5420-5427.
- [7] HE Y J, LI W. A lightweight deep learning approach for encrypted proxy traffic detection[J]. Security and Communication Networks, 2025, 2025(1): 4469975.
- [8] XU H, CHENG Z, LI S, et al. ProxyKiller: an anonymous proxy traffic attack model based on traffic behavior graphs[C]//Computer Security—ESORICS 2024. Bydgoszcz, Poland, 2024-09-16—2024-09-20. Cham: Springer, 2024: 162-181.
- [9] JIA H, CHEN Y, TANG Z. SPTC: signature-based cross-protocol encrypted proxy traffic classification approach[C]//Information and Communications Security: ICICS 2025. Nanjing, China, 2025-10-29—2025-10-31. Singapore: Springer, 2026: 457-475.
- [10] GU Z, GOU G, HOU C, et al. LFETT2021: a large-scale fine-grained encrypted tunnel traffic dataset[C]//Proceedings of the 2021 IEEE 20th International Conference on Trust, Security and Privacy in Computing and Communications. Shenyang, China, 2021-10-20—2021-10-22. Piscataway: IEEE Press, 2021: 240-249.
- [11] WANG T, XIE X, WANG W, et al. NetMamba: efficient network traffic classification via pre-training unidirectional Mamba[C]//Proceedings of the 2024 IEEE 32nd International Conference on Network Protocols. Charleroi, Belgium, 2024-10-28—2024-10-31. Piscataway: IEEE Press, 2024: 1-11.
- [12] DENG X W, ZHAO R J, ZHAN M W, et al. Multi-modal datagram representation with spatial-temporal state space models and inter-flow contrastive learning for encrypted traffic classification[C]//Information and Communications Security: ICICS 2025. Nanjing, China, 2025-10-29—2025-10-31. Singapore: Springer, 2026: 476-492.
- [13] 刘涛涛, 付钰, 俞艺涵, 等. 基于并行流量图和图神经网络的加密流量分类方法[J]. 通信学报, 2025, 46(6): 45-59.
- [14] ZHANG H Z, YUE H D, XIAO X, et al. Revolutionizing encrypted traffic classification with MH-Net: a multi-view heterogeneous graph model[C]//Proceedings of the 39th AAAI Conference on Artificial Intelligence. Philadelphia, PA, USA, 2025-02-25—2025-03-04. Palo Alto: AAAI Press, 2025, 39(1): 1048-1056.
- [15] WEN J G, ZHANG X Y, XIE K, et al. Graph-based contrastive learning and clustering for open-world encrypted traffic classification[J]. IEEE Transactions on Information Forensics and Security, 2026, 21: 6363-6377.
- [16] LIU L, WEI Z L, LIU Z T, et al. Towards practical few-shot multi-tab website fingerprinting[C]//Proceedings of the 35th USENIX Security Symposium. Baltimore, MD, USA, 2026-08-12—2026-08-14. Berkeley: USENIX Association, 2026: 3083-3103.
- [17] WU M, SIPPE J, SIVAKUMAR D, et al. How the Great Firewall of China detects and blocks fully encrypted traffic[C]//Proceedings of the 32nd USENIX Security Symposium. Anaheim, CA, USA, 2023-08-09—2023-08-11. Berkeley: USENIX Association, 2023: 2653-2670.
- [18] 武雯欣, 徐国天, 朱广锐. 面向新型 V2Ray 类加密代理协议伪装变种流量的识别分类技术[J/OL]. 计算机工程, 2026[2026-09-05]. <https://doi.org/10.19678/j.issn.1000-3428.0253298>.
- [19] XUE D, STANLEY R, KUMAR P, et al. The discriminative power of cross-layer RTTs in fingerprinting proxy traffic[C]//Proceedings of the 2025 Network and Distributed System Security Symposium. San Diego, CA, USA, 2025-02-24—2025-02-28. Reston: Internet Society, 2025.
- [20] CERASUOLO F, NASCITA A, BOVENZI G, et al. MEMENTO: a novel approach for class incremental learning of encrypted traffic[J]. Computer Networks, 2024, 245: 110374.
- [21] WANG X M, WANG Y P, LAI Y X, et al. Reliable open-set network traffic classification[J]. IEEE Transactions on Information Forensics and Security, 2025, 20: 2313-2328.
- [22] ZHAO Y, DETTORI G, BOFFA M, et al. The sweet danger of sugar: debunking representation learning for encrypted traffic classification[C]//Proceedings of the ACM SIGCOMM 2025 Conference. Coimbra, Portugal, 2025-09-08—2025-09-11. New York: ACM Press, 2025: 296-310.
- [23] WICKRAMASINGHE N, SHAGHAGHI A, TSUDI K, et al. SoK: decoding the enigma of encrypted network traffic classifiers[C]//Proceedings of the 2025 IEEE Symposium on Security and Privacy. San Francisco, CA, USA, 2025-05-12—2025-05-15. Piscataway: IEEE Press, 2025: 1825-1843.
- [24] DING R Y, SUN L, DING Z Y, et al. Towards targeted and universal adversarial attacks against network traffic classification[J]. Computers & Security, 2025, 155: 104470.
- [25] LIU Z X, ZHAO Y, LIU Z T, et al. A hard-label black-box evasion attack against ML-based malicious traffic detection systems[C]//Proceedings of the 2026 Network and Distributed System Security Symposium. San Diego, CA, USA, 2026-02-23—2026-02-27. Reston: Internet Society, 2026.
- [26] ALI T, YOOSUF S, RABHI M, et al. Beyond RTT: an adversarially robust two-tiered approach for residential proxy detection[C]//Proceedings of the 2026 Network and Distributed System Security Symposium.

- sium. San Diego, CA, USA, 2026-02-23—2026-02-27. Reston: Internet Society, 2026.
- [27] XIE G R, LI Q, SHI Z N, et al. Defending against traffic analysis attacks with flexible in-network obfuscation[C]//Proceedings of the 23rd USENIX Symposium on Networked Systems Design and Implementation. Renton, WA, USA, 2026-05-04—2026-05-06. Berkeley: USENIX Association, 2026: 2043-2063.
- [28] GU A, DAO T. Mamba: linear-time sequence modeling with selective state spaces[C]//Proceedings of the First Conference on Language Modeling. Philadelphia, PA, USA, 2024-10-07—2024-10-09. [S.l.]: OpenReview, 2024.
- [29] KHOSLA P, TETERWAK P, WANG C, et al. Supervised contrastive learning[C]//Advances in Neural Information Processing Systems 33. Virtual Event, 2020-12-06—2020-12-12. Red Hook: Curran Associates, 2020: 18661-18673.
- [30] WANG T, XIE X, WANG W, et al. NetMamba+: a framework of pre-trained models for efficient and accurate network traffic classification [EB/OL]. arXiv: 2601.21792, 2026[2026-09-05]. <https://arxiv.org/abs/2601.21792>.
- [31] GRETTON A, BORGWARDT K M, RASCH M J, et al. A kernel two-sample test[J]. Journal of Machine Learning Research, 2012, 13(25): 723-773.

毛蔚轩, 于 2017
学与工程专业获
计算机网络应急技
师。研究领域为
趣包括: 入侵检
评估。Email: m



大学控制科
见任国家计
高级工程
。研究兴
漏洞挖掘
t.org.cn。