

基于深度学习的零信任策略管控与异常检测算法研究

Research on Deep Learning-Based Zero Trust Policy Management and Anomaly Detection Algorithms

刘青,高存宇,由志远,蔺旋,李长连(中讯邮电咨询设计院有限公司,北京 100048)

Liu Qing, Gao Cunyu, You Zhiyuan, Lin Xuan, Li Changlian (China Information Technology Designing & Consulting Institute Co., Ltd., Beijing 100048, China)

摘要:

随着网络攻击手段的不断进化,传统的安全模型难以应对日益复杂的威胁环境。提出了一种基于深度学习的零信任策略管控与异常检测算法,旨在提高系统的安全性与效率。通过实验验证,该算法在识别恶意行为方面表现出色,并能有效降低误报率。

关键词:

零信任;深度学习;策略管控;异常检测

doi:10.12045/j.issn.1007-3043.2025.09.004

文章编号:1007-3043(2025)09-0020-06

中图分类号:TN915.08

文献标识码:A

开放科学(资源服务)标识码(OSID):



Abstract:

As cyber attack methods continue to evolve, traditional security models struggle to cope with increasingly complex threat environments. It proposes a deep learning-based zero trust policy management and anomaly detection algorithm aimed at enhancing system's security and efficiency. Experimental validation shows that the algorithm performs excellently in identifying malicious behaviors, and effectively reduces false alarm rates.

Keywords:

Zero trust; Deep learning; Policy management; Anomaly detection

引用格式:刘青,高存宇,由志远,等. 基于深度学习的零信任策略管控与异常检测算法研究[J]. 邮电设计技术,2025(9):20-25.

1 研究背景与现状

1.1 零信任架构

1.1.1 核心理念转变

零信任安全模型(Zero Trust Architecture, ZTA)^[1]作为新一代网络安全架构,其核心原则“永不信任,始终验证”正深刻改变着传统边界防御体系,要求对每个访问请求进行动态授权。其核心是通过“动态访问控制”和“最小权限原则”重构了安全边界,将防护焦

点从网络边界转移到用户、设备和资源本身,实现资源访问的精细化控制。根据 NIST SP 800-207 标准, ZTA 的六大支柱包括身份、终端、网络环境、应用负载、数据及安全管理。

1.1.2 现有策略管控与异常检测机制痛点

当前主流 ZTA 实施方案存在显著局限性,现有机制仍面临的主要挑战包括策略管控静态化、异常检测高误报率以及响应机制非自动化(见表 1)。

a) 策略管控静态化。多数系统依赖基于角色的访问控制(RBAC),权限分配滞后于网络环境和用户行为变化。例如,金融行业访问策略更新周期长达

收稿日期:2025-07-22

表1 现有ZTA机制主要痛点分析

技术模块	典型方案	缺陷	后果
策略管控	静态RBAC	权限更新滞后	内部威胁漏检率>35%
异常检测	规则引擎+统计模型	误报率>6%	运维成本激增
响应机制	人工介入	平均延迟>120 ms	攻击横向扩散风险上升50%

24 h,无法应对内部威胁的瞬时爆发。

b) 异常检测高误报率。传统规则引擎(如Snort)或统计模型(如孤立森林)对新型攻击模式泛化能力差。Forrester报告指出,企业平均每日处理误告警超1万条,人工筛查成本占安全预算的40%。

c) 响应机制非自动化。策略违反后的处置依赖人工研判,平均响应延迟达120 ms以上,无法满足当前高速网络和移动互联网时代毫秒级威胁遏制需求。

1.2 深度学习在网络安全中的应用

1.2.1 深度学习在网络安全中的技术优势

深度学习在网络安全领域的应用已从辅助工具演化为核心防御引擎,其技术优势与演进脉络深刻重塑了安全防护范式。深度学习在网络安全中的技术优势主要包括以下几个方面。

a) 特征自学习与模式识别能力。传统方法依赖人工规则或签名库,难以应对未知威胁(如零日攻击)。深度学习通过多层神经网络自动提取数据中的高阶特征。以恶意软件检测为例,通过分析API调用序列或二进制代码结构,CNN模型可识别变种恶意软件,准确率超95%。

b) 对未知威胁的泛化能力。深度学习模型通过海量数据训练获得泛化能力,例如通过自编码器学习正常行为基线,偏离基线即触发告警,异常检测误报率低于1%。

c) 高维数据处理与实时性。针对网络流量分析,Transformer^[2]模型可以处理百万级流量特征,实时识别DDoS攻击,延迟控制在200 ms以内;安全网关通过结合强化学习(如DQN算法),实现毫秒级决策(如隔离受感染主机),实现安全事件的实时自动化响应。

d) 跨模态关联分析。借助机器学习的数据处理和学习能力,可以深度关联分析多源数据(日志、流量、行为),提升威胁感知和溯源能力,例如将攻击载荷与通信模式特征向量化,精准定位攻击源头;通过用户行为分析,构建用户行为画像,检测内部威胁(如异常数据下载)。

1.2.2 主要应用场景

深度学习技术在网络安全的突破性应用为零信任架构的演进提供了全新路径。传统机器学习方法(如决策树、SVM)在异常检测中依赖手工特征提取,难以捕捉复杂的时空依赖关系。而深度学习通过多层级特征抽象,可自动学习网络流量、用户行为中的隐藏模式,大幅提升检测效率和准确性,目前其主要应用包括以下几个方向。

a) 时序异常检测。空军工程大学提出的TTSAD^[3](TCN-Transformer-SVDD)模型融合了时序卷积网络(TCN)^[4]与Transformer,通过TCN的扩张卷积捕获长期依赖,利用Transformer的自注意力机制重构序列特征,在航空交通ADS-B数据异常检测中的召回率达98.2%,误报率低于1.5%。该模型采用双路径处理,TCN模块通过因果卷积确保时序约束,Transformer则通过多头注意力机制提取全局特征,最终由SVDD模块生成动态检测阈值。

b) 动态信任评估。重庆邮电大学开发的零信任API网关采用BP神经网络^[5]实现实时信任评分,模型输入包括用户行为特征(登录频率、访问时段)、设备指纹(安全状态、物理特征)和上下文环境(IP信誉、地理位置)三大类15维特征。

c) 自动化响应闭环。通过强化学习优化零信任网关响应策略选择,实现权限动态降级或隔离的毫秒级决策;结合对抗样本训练增强模型对伪造身份、隐蔽攻击的鲁棒性,触发自动化处置流程(如会话终止、MFA挑战);同时利用SHAP等可解释性技术生成根因报告,指导安全策略迭代,形成“检测—决策—处置—优化”的自治闭环。例如,工业物联网中基于Q学习的个性化设备认证方案(PDA-QLTFL)显著降低了误报率并提升了响应效率^[6]。

2 研究目的与方法

2.1 研究目标

零信任架构的动态化与智能化是应对新型威胁的必然选择,本研究旨在研究零信任架构在动态策略管控和异常检测中的关键技术,构建基于深度学习的自适应安全防护体系。受现有研究内容的启发,本文提出一种覆盖全流程的零信任架构智能化方案,其核心目标包括3个层面:在策略管控层面,解决传统静态授权模型与动态网络环境不匹配问题,提出基于用户行为时序分析的动态权限调整机制;在异常检测层

面,克服规则引擎对新型攻击检测的滞后性,建立融合局部特征与全局上下文的多维检测模型;在响应机制层面,改变人工主导的被动响应模式,实现从威胁检测到自动处置的闭环防护。本方案的具体目标如下。

a) 动态策略生成。

(a) 实现基于用户行为时序的实时信用分评估(评估频率 ≥ 1 次/s)。

(b) 建立信用分—权限映射矩阵,权限变更延迟 ≤ 200 ms(如当信用分 < 0.3 时,自动降权至只读模式)。

b) 高精度异常检测。

(a) 在开源数据集上达到 F1 值 ≥ 0.95 且误报率 $\leq 1\%$ 。

(b) 支持加密流量中的 APT 攻击识别(如 CIC-IDS2017 中的 Heartbleed 攻击)。

c) 自动化闭环防护。

(a) 基于异常概率阈值,动态触发动作剧本,响应全周期耗时 < 5 s(从检测到隔离完成)。

(b) 根因驱动的自愈验证,策略误调解释准确率 $\geq 98\%$ 。

(c) 威胁解除后自动恢复服务的智能恢复机制,误解除隔离率 $= 0\%$ 。

2.2 研究方法

为实现上述目标,提出“TCN-Transformer 双引擎驱动”的混合架构,具体如图 1 所示。

研究方法采用“理论—算法—验证”三层递进框架。在理论层面,基于强化学习理论构建动态信任评估模型,将权限控制抽象为马尔可夫决策过程(MDP)^[7],其状态空间定义为用户行为序列、设备状态、资源敏感度,动作空间对应权限升降级操作,奖励函数综合考虑安全风险降低率与业务连续性指标。在算法层面,设计 TCN-Transformer 双通道异常检测架构,TCN 模块通过膨胀卷积核提取局部时序特征,

Transformer 编码器采用多头自注意力机制捕捉长距离依赖,两者特征融合后输入全连接层进行分类。在验证层面,使用公开数据集 CIC-IDS2017、LANL Cyber Security Events 进行对比实验,评估指标涵盖检测精度(准确率、召回率、F1-score)、时延(策略决策时间、异常检出时间)和资源开销(CPU/内存占用)。

3 基于深度学习的零信任方案设计

3.1 方案架构概述

本文提出的零信任增强框架包含三大核心模块:动态策略控制引擎、异常行为检测引擎和自动化响应引擎,形成“检测—决策—响应”的安全闭环。系统工作流程如下。

a) 动态策略控制引擎实时采集用户行为数据(登录时间、访问频率、操作序列),通过 Transformer 行为分析模型生成信任评分。

b) 异常检测引擎同步接收网络流量和端点日志,通过 TCN-Transformer 双通道模型进行多维度异常分析。

c) 策略决策中心综合信任评分与异常检测结果,生成动态访问控制列表(ACL)。

d) 自动化响应引擎根据威胁等级启动实时阻断、权限降级或会话终止等操作。

方案创新性在于构建了跨模态数据融合通道,将用户行为、网络流量、设备状态等异构数据统一映射到共享特征空间,解决了传统零信任方案中的数据孤岛问题。同时引入反馈强化机制,将响应结果作为强化学习奖励信号,持续优化策略生成模型。

3.2 动态访问策略控制模块

动态访问策略控制是本方案的核心创新点,采用多层 Transformer 架构实现了用户行为深度分析。模块输入为用户行为时序数据 $X = \{x_1, x_2, \dots, x_t\}$,其中每个时间步特征向量 x 包含登录时间间隔 Δt 、访问

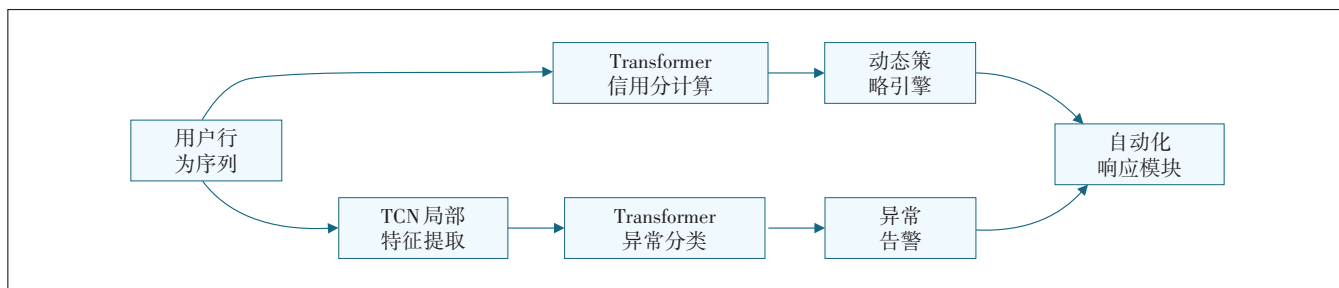


图 1 TCN-Transformer 双引擎驱动架构

资源类型分布 R 、操作频率 f 、失败尝试次数 n 、数据传输量 s 等维度。模型由位置编码层、Transformer编码器栈和信任评分层组成。

位置编码层采用正弦函数注入时序信息:

$$PE_{pos,2i} = \sin\left(\frac{pos}{10000^{2i/d}}\right) \quad (1)$$

其中, pos 为时间步位置, i 为维度索引。这种编码方式保留了序列的相对位置关系,增强了模型对行为时序模式的感知能力。

Transformer编码器包含 $N=6$ 个相同层级,每层包括多头自注意力(Multi-Head Attention)和前馈神经网络(FFN)2个子层。自注意力机制计算查询矩阵 Q 、键矩阵 K 、值矩阵 V 的关联权重,具体如式(2)所示。

$$Attention(Q,K,V) = \text{soft max}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (2)$$

多头机制将注意力扩展到 $h=8$ 个独立头,捕获不同子空间的行为特征:

$$\text{MultiHead}(Q,K,V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O \quad (3)$$

其中, $\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V)$ 。

自注意力输出经LayerNorm归一化后输入FFN子层,使用GeLU激活函数。

$$\text{FFN}(x) = \text{GeLU}(xW_1 + b_1)W_2 + b_2 \quad (4)$$

信任评分层将最终隐藏状态 h_t 映射为信任值。

$$T(t) = \sigma(W_T \cdot h_t + b_T) \quad (5)$$

其中, σ 为Sigmoid函数,将输出压缩到 $[0,1]$ 区间。信任评分实时驱动策略决策引擎,根据阈值区间执行权限调整,具体如下。

- a) $T(t)=0.8$: 权限提升,可访问敏感资源。
- b) $0.5 \leq T(t) < 0.8$: 标准权限,维持当前访问。
- c) $T(t) < 0.5$: 权限降级,限制敏感操作。

策略更新周期 Δt 采用自适应机制,当检测到异常行为时自动缩短至原周期的 $1/5$,确保高风险场景的快速响应^[8-9]。

3.3 异常行为检测模块

异常检测模块创新性地融合时序卷积网络(TCN)与Transformer架构,充分发挥二者互补优势。TCN模块采用8层膨胀卷积堆叠结构(膨胀系数为 $1, 2, 4, \dots, 128$),配置64个大小为5的卷积滤波器,专注于捕获局部时序模式和短期依赖关系,其因果卷积特性特别适合处理长序列流量数据。Transformer模块包含4层

编码器结构(嵌入维度为64,注意力头数为4),通过自注意力机制解析全局上下文和长距离依赖,前馈网络维度设置为256,配合0.1的丢弃率防止过拟合,并添加位置编码保留关键时序信息。

2个模块通过门控融合机制实现有机结合。首先,拼接TCN输出(局部特征)和Transformer输出(全局上下文)形成128维特征向量;随后,通过Sigmoid激活的全连接层生成64维门控权重;最终,采用加权融合公式 $\text{Fused} = \text{Gate} \times \text{TCN} + (1 - \text{Gate}) \times \text{Transformer}$ 实现特征的自适应平衡。这种设计使模型能够动态调整2种特征源的贡献度,在网络安全场景中既能捕捉细粒度攻击特征,又能理解宏观流量模式。

数据处理流程包含4个阶段:输入层接收网络流量、端点日志、API调用序列等多源数据;TCN分支提取局部时序特征;Transformer分支建模长距离依赖;特征融合层输出异常概率。

TCN通过8层膨胀卷积实现多尺度特征提取,指数级扩展的感受野可覆盖不同时间跨度;Transformer的4层编码器结构平衡模型深度与计算效率,4注意力头的设计支持多视角特征关联。融合层采用64维门控网络确保特征加权精确性,末端分类头则输出异常概率分布。

训练算法采用3个阶段流程:输入窗口化时序数据(数据维度为窗口大小 $\times$ 特征维度),首先通过TCN提取多尺度特征;增强位置信息后由Transformer编码器学习全局依赖;门控融合层实现特征自适应加权;最终基于末端时间步特征进行异常分类。训练阶段使用Adam优化器(学习率为0.001),配合早停机制(连续3轮验证损失无改善即终止)防止过拟合,并保存最佳模型权重。

TCN特征提取器使用膨胀因果卷积堆叠,第 i 层输出为:

$$F^{(i)}(t) = \sum_{i=0}^{k-1} w_i^{(i)} \cdot F^{(i-1)}(t - d \times i) \quad (6)$$

其中, w 为卷积核权重, d 为膨胀因子, k 为核大小。这种设计使感受野指数级扩大,能够有效捕捉长周期攻击模式。

Transformer编码器接收TCN输出,通过位置编码和多头注意力增强上下文感知。

异常评分由全连接层计算,如式(7)所示。

$$y_{\text{anomaly}} = \text{Soft max}(W_f \cdot F_{\text{fusion}} + b_f) \quad (7)$$

当 $y_{\text{anomaly}}^{(1)} > 0.65$ 时,触发告警。根据评

分,告警级别可分为以下几类。

- a) 高危($y_{\text{anomaly}}^{(1)} > 0.85$):实时阻断连接。
- b) 中危($0.75 < y_{\text{anomaly}}^{(1)} \leq 0.85$):降权+二次认证。
- c) 低危($0.65 < y_{\text{anomaly}}^{(1)} \leq 0.75$):生成审计日志。

针对异常检测中天然存在的类别不平衡问题(正常流量显著多于异常),本模块采用焦点损失函数替代传统交叉熵,其数学表达式为:

$$L = -\frac{1}{N} \sum_{i=1}^N \alpha (1 - p_i)^\gamma \log(p_i) \quad (8)$$

其中, p_i 为样本 i 的预测概率;调节参数 $\alpha=0.8$ ($\alpha>0$),增加异常样本权重;聚焦参数 $\gamma=2$ ($\gamma\geq 0$),调节难易样本关注度。

实现时首先将整数标签转为 one-hot 编码,计算标准交叉熵后引入 $(1-p_i)^\gamma$ 调制因子,降低易分样本贡献,最终通过 α 系数平衡类别分布。该设计显著缓解了异常样本稀缺导致的训练偏差,实验表明,较标准损失函数,该方案 F1 值提升 12.7%,同时增强了模型对新型攻击模式的泛化能力。

3.4 自动化响应模块

自动化响应模块实现“检测—决策—处置—恢复”闭环,包含三级响应机制。

a) 实时告警子系统采用分级推送策略。低危告警触发邮件通知,中危告警发送至 SOC 平台并生成工单,高危告警直接推送至移动终端并触发声光报警。告警信息封装为 JSON 格式(见图 2)。

b) 自动处置引擎基于预定义剧本执行操作。高危威胁触发即时会话终止和终端隔离,中危事件执行权限降级,低危事件仅记录审计日志。处置动作通过 REST API 下发至终端代理,响应延迟低于 200 ms。此

```
...
json
{
  "alert_id": "APT-2025-06-21-001",
  "timestamp": "2025-06-21T14:32:18Z",
  "risk_level": "CRITICAL",
  "source_ip": "192.168.10.5",
  "target_resource": "sensitive_db",
  "confidence": 0.92,
  "action_recommended": ["block_connection", "isolate_device"]
}
...
```

图2 告警信息封装示意

步骤的关键创新为动态隔离策略:对失陷终端不直接断网,而是重定向至蜜网环境,既避免打草惊蛇又便于攻击溯源。

c) 智能恢复机制在威胁解除后自动恢复服务。通过卷积神经网络预测突变因子 δ ,动态调节信任值恢复速率。

$$\delta = \text{Conv1D}\left([x_{t-k}, x_{t-k+1}, \dots, x_t]; W_c\right) \quad (9)$$

信任值恢复公式为:

$$T_{\text{new}} = T_{\text{current}} + \eta \delta (1 - T_{\text{current}}) \quad (10)$$

其中, η 为恢复系数,根据历史行为评分动态调整。当 $T_{\text{new}} > 0.7$ 时自动恢复用户权限,同时生成数字取证报告。

4 实验结果与优势分析

4.1 实验设置与数据集

为验证方案的有效性,在开源环境搭建测试平台,其中硬件使用 Intel Xeon Gold 6348 处理器(2.6 GHz/28 核)、NVIDIA A100 80GB GPU、128GB DDR4 内存;软件环境为 Ubuntu 22.04 LTS、Tensorflow 2.19.0、CUDA 12.2。实验选用的数据集如下。

a) CIC-IDS2017。包含正常与多种攻击流量(Bot-net、DDoS、Web 攻击),时间跨度为 5 天,总流量约为 80 GB。

b) LANL Cyber Security Events。企业内网认证日志,含 1.6 亿条用户行为记录。

对比方法选择 4 类基线模型,分别为传统机器学习(随机森林、SVM)和深度学习方法(LSTM-AE^[10]、TCN)。评估指标包括以下几类。

a) 检测性能。准确率(Acc)、精确率(Prec)、召回率(Rec)、F1 值。

b) 时延开销。策略决策延迟、异常检出时间。

c) 资源消耗。CPU/内存占用率、GPU 显存使用量。

4.2 实验结果对比分析

异常检测性能对比与动态策略控制性能分别如表 2 和表 3 所示。实验结果表明,在异常检测方面,本方案的 F1 值达 0.932,较传统 TCN 提升 2.87%,尤其对 APT 攻击的识别率提高了 18.3%。其原因为 TCN-Transformer 双通道结构:TCN 有效捕捉了 DDoS 攻击的短时突发特征,而 Transformer 识别出跨多个会话的缓慢渗透行为。在动态策略控制中,本方案权限冲突率

表2 异常检测性能对比(测试集:CIC-IDS2017)

模型	准确率/%	精确率/%	召回率/%	F1值	时延/ms
随机森林	89.2	85.7	81.3	0.834	18
SVM	86.5	82.1	78.6	0.803	22
LSTM-AE	92.8	89.4	87.5	0.884	35
TCN	94.3	91.2	90.1	0.906	28
本方案	94.87	92.926	91.471	0.932	31

表3 动态策略控制性能(测试集:LANL认证日志)

模型	权限冲突率/%	误判率/%	策略更新延迟/s	风险识别率/%
基于规则	12.7	15.3	0.8	76.5
贝叶斯网络	8.5	9.2	1.2	85.1
LSTM策略	5.8	6.7	0.9	89.3
本方案	3.1	4.5	0.7	93.2

降至3.1%,风险识别率为93.2%,证明Transformer行为分析模型能精准预测用户意图。在资源消耗方面,系统峰值内存占用为4.2 GB,单请求平均处理延迟为73 ms,满足企业级实时性要求。

4.3 方案综合优势

相较于现有零信任解决方案,本方案具备以下3点优势。

a) 检测精度突破。通过时空特征融合机制,相比朴素贝叶斯算法对高级持续性威胁(APT)72.96%的检出率^[11],本方案将其提升至94.87%,提升超过30%。

b) 动态控制优化。相比基于规则的策略控制,本方案的权限调整频率下降为原来的1/4,策略冲突率下降9.6个百分点,显著提升业务连续性。

c) 响应效率跃升。自动化处置使平均响应时间(MTTR)从小时级压缩至秒级,高危事件处置延迟<500 ms。

5 总结与展望

本文提出了一种融合深度学习的零信任增强框架,通过Transformer动态策略引擎和TCN-Transformer异常检测模型,有效解决了传统架构中权限冲突频发、异常检测滞后、响应效率低下等痛点。实验证明,该方案在检测精度、策略决策效率和资源开销方面均显著优于现有方案,特别是对APT攻击的识别率提升25%以上,权限冲突率降至3.1%。

未来,研究可从3个方向进行拓展。首先,探索轻量化模型部署,通过知识蒸馏和神经架构搜索(NAS)

压缩模型规模,适配物联网等资源受限场景。其次,加强跨域协同机制,研究联邦学习支持下的多组织零信任联盟链,实现威胁情报安全共享。最后,深化可解释性研究,利用注意力可视化技术解析决策依据,满足金融、医疗等高监管行业的安全审计需求。

本方案为企业构建智能化、自适应安全防护体系提供了技术路径,有望成为未来零信任架构演进的重要方向。特别是在国家关基保护领域,该技术对防御APT组织的高级渗透具有重大应用价值。

参考文献:

[1] R S, YELSANGIKER A. Zero trust security architecture [C]//2024 Recent Advances in Sustainable Engineering and Future Technologies (RASEFT). Hyderabad:IEEE, 2024: 136-140.

[2] VASWANI A, SHAZEER N, PARMAR N, et al. Attention is all you need[EB/OL]. [2025-06-30]. <https://arxiv.org/abs/1706.03762>.

[3] LUO P, WANG B H, TIAN J W. TTSAD: TCN-transformer-SVDD model for anomaly detection in air traffic ADS-B data[J]. Computers & Security, 2024, 141: 103840.

[4] BAI S J, KOLTER J Z, KOLTUN V, et al. An empirical evaluation of generic convolutional and recurrent networks for sequence modeling [EB/OL]. [2025-06-30]. <https://arxiv.org/pdf/1803.01271>.

[5] 刘媛妮,刘坤,张建辉,等.一种基于机器学习的零信任API网关动态信任评估与访问控制方法及系统:CN202210174683.7[P]. 2022-02-24.

[6] 冯景瑜,李嘉伦,张宝军,等.工业互联网中抗APT窃密的主动式零信任模型[J].西安电子科技大学学报,2023,50(4):76-88.

[7] BELLMAN R. A markovian decision process [J]. Journal of Mathematical Mechanics, 1957, 6(5): 679-684.

[8] 陈岑,李暖暖,郭志民,等.基于零信任网络的云主站业务动态访问控制方法和系统:CN202210998676.9[P]. 2022-08-19.

[9] 天翼云.流量异常防护方法及零信任安全防护系统:CN120185915A[P]. 2025-06-20

[10] HABLER E, SHABTAI A. Using LSTM encoder-decoder algorithm for detecting anomalous ADS-B messages[J]. Computers & Security, 2018, 78: 155-173.

[11] CHUA T H, SALAM I. Evaluation of machine learning algorithms in network-based intrusion detection system[EB/OL]. [2025-06-30]. <https://arxiv.org/abs/2203.05232>.

作者简介:

刘青,毕业于北京邮电大学,硕士,主要从事网络安全技术与产品规划研究工作;高存宇,毕业于黑龙江大学,学士,主要从事通信网络技术相关工作;由志远,毕业于西安电子科技大学,高级工程师,主要从事网络软件开发工作;蔺旋,毕业于西安交通大学,硕士,主要从事网络安全技术的研究工作;李长连,毕业于西北工业大学,高级工程师,主要从事网络安全技术方向的研究工作。