



专题：智能通信技术

信息网络内生恶意行为检测框架

涂哲, 周华春, 李坤, 王玮琳
(北京交通大学, 北京 100044)

摘要: “内生安全”赋予信息网络自学习、自成长的能力, 是构建可信智能通信网络不可或缺的重要组成部分。面向信息网络“内生安全”, 提出了一种内生恶意行为检测框架, 变被动防御为主动拦截。同时, 对内生恶意行为检测框架中五大关键组件进行了建模分析, 并对自学习、自成长的恶意行为检测机制进行了阐述。最后, 搭建原型系统并进行了实验, 实验结果表明了检测框架的可行性和有效性。

关键词: 信息网络; 内生安全; 恶意行为检测; 安全框架

中图分类号: TP393

文献标识码: A

doi: 10.11959/j.issn.1000-0801.2020282

Endogenous detection framework of malicious behavior in information network

TU Zhe, ZHOU Huachun, LI Kun, WANG Weilin
Beijing Jiaotong University, Beijing 100044, China

Abstract: “Endogenous security” gives the information network the ability of self-learning and self-growth, and it’s the important part of building a trusted intelligent communication network. “Endogenous security” was combined with malicious behavior detection, and an endogenous detection framework of malicious behavior with endogenous security capabilities was proposed. Passive defense was turned into active interception. In addition, the five key components in the endogenous detection framework were also modeled and analyzed, and explained the self-learning and self-growth malicious behavior detection mechanism. Finally, the deployment method of the prototype system was introduced. Preliminary test results show the proposed framework is feasible and effective.

Key words: information network, endogenous security, malicious behavior detection, security architecture

1 引言

IMPREVA 威胁研究实验室于 2020 年发布的

最新研究报告表明, 恶意的机器流量在全网流量中占比逐年上升, 现已升至有史以来的最高比例 24.1%^[1]。恶意流量已成为威胁信息网络安全、制

收稿日期: 2020-09-02; 修回日期: 2020-10-10

基金项目: 国家重点研发计划项目 (No.2018YFA0701604); 国家自然科学基金资助项目 (No.61802014, No.U1530118); 国家高技术研究发展计划 (“863” 计划) 基金资助项目 (No.2015AA015702)

Foundation Items: The National Key Research and Development Program of China (No.2018YFA0701604), The National Natural Science Foundation of China (No.61802014, No.U1530118), The National High Technology of China (863 Program) (No.2015AA015702)



约网络传输效率的重大障碍，有效地抑制恶意流量对扩充网络容量起着至关重要的作用。然而，信息网络在设计初始缺乏安全性的考量，传统信息安全通过依附于专业硬件的“外挂式”安全保护机制和“补丁式”的安全漏洞修复方式得以保障。由于缺乏从底层网络到上层应用的全局统筹，传统信息网络难以从根本上解决层出不穷的网络安全问题，无法有效地应对动态变化的信息网络环境^[2]。因此，迫切需要研究出一种新的网络安全技术体系。

“内生安全”突破传统思路，变革网络安全范式，是一种全新的网络安全体系。通过将安全能力与网络设备深度融合，网络安全不再依附于特定的安全设备与安全机制，信息网络因此拥有强健的安全基因，具备由内而外不断生长的安全能力。“内生安全”具有自适应、自主、自生长等特征，是构建安全可信智能通信网络不可或缺的重要组成部分^[3]。

近些年来，许多专家学者对“内生安全”架构进行了积极的研究。参考文献[4]分析了 5G 网络中的安全形势以及内生安全新需求，并对 5G 内生安全架构发展趋势和关键技术进行了阐述。为了应对 6G 网络中的安全挑战，参考文献[3]提出了 6G 网络内生安全体系结构和相关技术要求。此外，还有不少研究聚焦于“内生安全”的具体技术和应用场景。参考文献[5]提出了一种由身份管

理者、本地 DHCP 服务器、ID 验证者、边界路由器和审计代理 5 种特定安全功能组件组成的具有内生安全特性的网络架构 NAIS。参考文献[6]从具有随机性、时变性和互易性等特征的无线信道展开研究，提出了一种基于射频指纹和信道密钥，融合认证、加密、传输为一体的内生安全通信框架。在安全光通信方面，参考文献[7]提出了一种“通密一体，防传融合”的内生安全光通信结构，并给出了基于“微元”的“三区耦合，两路简并”的微元内生安全模型。

目前，关于“内生安全”的研究大都围绕着内生安全网络架构和特定场景下的技术应用，尚未有针对内生恶意行为检测方面的具体研究。恶意通信行为具有复杂性高、不确定性大、隐蔽性强等特点，是构建安全可信通信网络的重大障碍。因此，本文围绕恶意行为检测展开，以减少信息网络恶意通信行为为出发点，从“内生安全”维度建立“用户身份-通信行为-用户信誉-安全控制”连接关系，基于行为知识库和分布式检测方法提出一种具有内生安全的恶意通信行为检测框架，旨在赋予通信网络内生恶意行为检测的能力。

2 内生恶意行为检测框架

内生恶意行为检测框架如图 1 所示。内生恶意行为检测框架突破传统思路，是一种全新的网络安全范式。通过在边界路由器中部署恶意行为

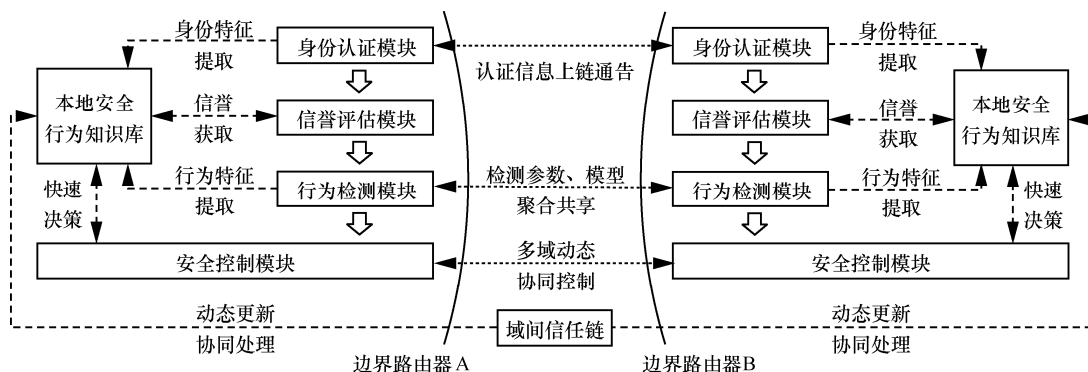


图 1 内生恶意行为检测框架

检测代理,将内生安全能力与网络硬件深度融合,变革传统信息网络“外挂式”安全保护机制;通过分布式感知恶意行为、聚合共享参数和构建域间安全行为知识库动态更新、协同处理的信任链的方法,赋予网络应对恶意通信行为时的自学习、自成长能力,变更传统信息网络“补丁式”安全漏洞修复方式。

为了构建具有“内生安全”的恶意行为检测框架,内生恶意行为检测框架提供了 5 个重要的安全功能组件:身份认证模块、信誉评估模块、行为检测模块、安全控制模块和本地安全行为知识库。身份认证模块负责对接入设备身份进行认证,通过构建基于身份的可信协议,满足差异化用户身份认证需求,及时快速地将恶意用户阻隔;信誉评估模块对身份认证后的用户行为进行评估,建立基于用户历史通信行为的信誉评估体系,动态分析通信伙伴与链接之间的信任关系,实现细粒度身份评估;行为检测模块负责对信息网络中的恶意行为进行检测,从通信行为的角度,根据用户及其通信行为的链接关系对恶意流量进行分析、鉴别;安全控制模块是整个框架中负责动作执行的组件,能够根据行为检测结果实时地对网络接入进行管控,实现对用户不同粒度的访问控制与行为约束。此外,安全控制模块也能够根据本地安全行为知识库的泛化推理能力,对恶意行为进行快速决策,阻断恶意用户破坏行为;本地安全行为知识库是内生恶意行为检测框架中的关键组件,通过提取身份、信誉及行为等特征,建立“用户身份-通信行为-用户信誉-安全控制”链接关系,增强网络的感知能力、推理能力和决策能力。

内生恶意行为检测框架通过“本地检测、多域协同”赋予信息网络“内生安全”的能力。一方面,通过域内五大安全组件的通力协作,形成了基于身份与行为的认证评估体系,构建了具有“用户身份-通信行为-用户信誉-安全控制”连接

关系的本地安全行为知识库,实现了基于用户行为的检测、管控,有效地对本地恶意通信行为进行阻断。另一方面,采用身份认证信息上链通告、恶意行为检测模型参数聚合共享、安全控制多域协同以及域间信任链动态更新等方法,形成了身份认证、行为检测、安全管控和知识库共享等层面的多域动态更新、协同处理。

3 内生恶意行为检测理论模型

围绕内生恶意行为检测框架中的几个关键组件:身份认证模块、信誉评估模块、行为检测模块、安全控制模块和安全行为知识库进行分析,探讨各组件的基本原理,对模块功能进行系统性建模。

3.1 身份认证模块

身份认证模块负责对接入设备、用户身份进行认证,是内生恶意行为检测框架中第一道安全保护屏障。为了全面地表征接入用户身份,身份认证模块对接入用户身份进行统一标识。用户身份标识 UID 可以定义如下:

$$UID \triangleq \Omega(ID_{User}, ID_{Device}) \quad (1)$$

式(1)中 $\Omega()$ 为标识生成函数,而 ID_{User} 、 ID_{Device} 分别表征用户属性集合和接入设备属性集合:

$$ID_{User} \triangleq \{User_N, User_C, User_p \dots\} \quad (2)$$

$$ID_{Device} \triangleq \{Device_N, Device_{SN}, Device_{PM}, Device_{MAC} \dots\} \quad (3)$$

ID_{User} 集合包括用户名 $User_N$ 、用户族群 $User_C$ 、用户初始权限 $User_p$ 等关于用户属性的特征;接入 ID_{Device} 集合则包括设备名称 $Device_N$ 、设备序列号 $Device_{SN}$ 、设备型号 $Device_{PM}$ 、设备物理地址 $Device_{MAC}$ 等表征设备属性的特征。

对于同一个边界路由器,在某个时刻 t 内可能存在 n 个用户接入,因此由式(1)可以将边界路由器在时刻 t 的用户身份标识集合 UID_n^t 写出,



如式(4)所示,其中 UID_i 表示第 i 个接入用户身份标识。

$$UID'_n \triangleq \{UID_1, UID_2, \dots, UID_i, \dots, UID_n\} \quad (4)$$

为了更好地对身份认证模块进行建模,将身份认证结果集合表征为用户身份标识集合 UID'_n 映射的形式,见式(5)。其中 $Auth'_i$ 表示 t 时刻第 i 个接入用户的身份认证结果, $\theta()$ 为反映用户身份标识与认证结果间关系的映射。

$$\begin{bmatrix} Auth'_1 \\ \dots \\ Auth'_i \\ \dots \\ Auth'_n \end{bmatrix} \triangleq \theta(UID'_n) \triangleq \theta \left(\begin{bmatrix} UID_1 \\ \dots \\ UID_i \\ \dots \\ UID_n \end{bmatrix} \right) \quad (5)$$

3.2 信誉评估模块

信誉评估是整个内生恶意行为检测框架中的重要组件。信誉评估模块通过计算、分析用户的历史行为,能够准确地对用户信誉进行评估。为了更好地对信誉评估模块进行建模分析,首要的工作便是对用户历史行为进行了定义,将在一段时间内的用户历史行为集合 $Hist_Behav$ 表征如下:

$$Hist_Behav \triangleq \begin{bmatrix} AD, AL, TC, FC, \dots, L \\ SN, RN, NE, NR, \dots, C \\ SR, DR, SV, DV, \dots, A \end{bmatrix} \quad (6)$$

其中, L 、 C 、 A 分别对应着历史链接行为、历史通信行为和历史访问行为。历史链接行为包括平均链接建立时长 AD 、平均链接建立时延 AL 、链接建立总次数 TC 和链接建立失败次数 FC 等;历史通信行为包括发送的总数据量 SN 、接收的总数据量 RN 、发送错误数据量 NE 和重传分组个数 NR 等;历史访问行为包括安全资源请求数量 SR 、危险资源请求数量 DR 、安全域访问次数 SV 和危险域访问次数 DV 等。

为了使信誉评估更为合理,需要充分评估用户历史行为,在 t 时刻前,选取时间跨度为 T 的第 i 个用户的历史行为集合 $Hist_Behav_i^T$ 进行信誉

评估。 $Hist_Behav_i^T$ 可以表征如式(7)所示,其中 $Hist_Behav_j$ 表示在跨度为 T 的时间内第 j 个时间段的用户历史行为集合。

$$Hist_Behav_i^T \triangleq \{Hist_Behav_1, Hist_Behav_2, \dots, Hist_Behav_j, \dots\} \quad (7)$$

对于同一边界路由器中的 n 个接入用户,用户信誉是由历史行为和身份认证结果共同决定的,将用户信誉评估结果映射为用户历史行为 $Hist_Behav_i^T$ 与身份认证结果 $Auth'_i$ 的集合,见式(8)。其中 $Cred'_i$ 表示 t 时刻第 i 个用户的信誉评估, $\mu()$ 为表征用户信誉评估结果的计算函数。

$$\begin{bmatrix} Cred'_1 \\ \dots \\ Cred'_i \\ \dots \\ Cred'_n \end{bmatrix} \triangleq \mu \left(\begin{bmatrix} Hist_Behav_1^T \\ \dots \\ Hist_Behav_i^T \\ \dots \\ Hist_Behav_n^T \end{bmatrix}, \begin{bmatrix} Auth'_1 \\ \dots \\ Auth'_i \\ \dots \\ Auth'_n \end{bmatrix} \right) \quad (8)$$

3.3 行为检测模块

行为检测模块作为内生恶意行为检测框架中的核心组件,负责对信息网络中的恶意行为进行实时鉴别、分析,实现多类别恶意通信行为检测,为后续的安全控制模块提供可靠、细粒度的管控依据。实时通信行为特征 RT_Char 可以具体表征如式(9)。

$$RT_Char \triangleq \begin{bmatrix} FD, FP, FS, FN, FT, FF, \dots, Fc \\ LT, LD, LP, LL, LE, LS, \dots, Lc \\ SA, DA, SP, DP, SD, DD, \dots, Oc \end{bmatrix} \quad (9)$$

其中, Fc 、 Lc 、 Oc 分别表示流特征、链路特征和通信对象特征。流特征包括流持续时间 FD 、流协议号 FP 、流大小 FS 、流数量 FN 、流类型 FT 、流标志位集合 FF 等;链路特征包括链路类型 LT 、链路持续时间 LD 、链路协议类型 LP 、链路时延 LL 、链路加密方式 LE 、链路同步方式 LS 等;通信对象特征包括源地址 SA 、目的地址 DA 、源端口号 SP 、目的端口号 DP 、源域 SD 、目的域 DD 等。

在行为检测模块中, 用户行为的检测与分析不但与实时通信行为特征相关, 还与评估用户历史行为的信誉值相关。因此, 将用户行为检测结果表示如式 (10), 其中 Char_i^t 表示在同时拥有 n 个接入用户的边界路由器中, 第 i 个接入用户在 t 时刻的行为检测结果, RT_Char_i^t 为第 i 个接入用户的实时通信行为特征。 $\delta()$ 为行为检测函数, 常见的恶意行为检测方法有基于统计特性、基于机器学习算法和基于先验知识的检测方法。

$$\begin{bmatrix} \text{Char}_1^t \\ \dots \\ \text{Char}_i^t \\ \dots \\ \text{Char}_n^t \end{bmatrix} \triangleq \delta \left(\begin{bmatrix} \text{RT_Char}_1^t \\ \dots \\ \text{RT_Char}_i^t \\ \dots \\ \text{RT_Char}_n^t \end{bmatrix}, \begin{bmatrix} \text{Cred}_1^t \\ \dots \\ \text{Cred}_i^t \\ \dots \\ \text{Cred}_n^t \end{bmatrix} \right) \quad (10)$$

3.4 安全行为知识库

安全行为知识库是内生恶意行为检测框架中表征用户及其通信行为间的链接关系且具有强大关系推理能力的模块。通过对用户身份、用户信誉、用户行为特征等信息的知识抽取和知识存储, 安全行为知识库能构建基于“用户身份-通信行为-用户信誉-安全控制”的链接关系, 增强了信息网络深度感知、自我推理、快速决策的能力。

为了更好地对安全行为知识库进行建模, 首先需要对安全行为知识库中经过知识抽取后所存储的知识进行表征。可以用由用户身份、用户信誉、用户行为特征和安全控制结果等信息构成的集合表征行为知识库中存储的行为知识 Behav_K , 如式 (11) 所示, 其中 Ctrl 表征安全控制结果, 将在第 3.5 节中进行解释。

$$\text{Behav_K} \triangleq \{\text{UID}, \text{Hist_Behav}, \text{RT_Char}, \text{Ctrl} \dots\} \quad (11)$$

安全行为知识库由大量用户的行为知识组成的, 因此可以将安全行为知识库 BKD 表征如式 (12), 其中 Behav_K_i 表示安全行为知识库中所存储的第 i 个用户的行为知识, $\varphi()$ 表示用户行为之间表征链接关系的函数。

$$\text{BKD} \triangleq \varphi \left(\begin{bmatrix} \text{Behav_K}_1 \\ \dots \\ \text{Behav_K}_i \\ \dots \\ \text{Behav_K}_m \end{bmatrix} \right) \quad (12)$$

3.5 安全控制模块

安全控制模块是内生恶意行为检测框架中负责安全管控执行的关键组件。安全控制模块部署于行为检测模块之后, 能够根据不同用户的行为检测结果执行细粒度的安全控制。为了实现差异化安全管控, 安全控制模块通过对安全功能池 NFP 中的安全功能 NF 进行编排, 构造不同的安全功能链 NFC , 实现细粒度安全控制。

安全功能池 NFP 能够由式 (13) 所表征, 其中 $\text{NF}_1 \sim \text{NF}_m$ 表示域内所有的安全功能, 常见的安全功能包括深度分组检测、网络行为管理、蜜罐和防火墙等; $\text{NFI}_1 \sim \text{NFI}_m$ 表示所有安全功能所对应的安全功能信息, 包括性能、位置、使用情况等。

$$\text{NFP} \triangleq \begin{bmatrix} \text{NF}_1, \text{NFI}_1 \\ \text{NF}_2, \text{NFI}_2 \\ \dots \\ \text{NF}_m, \text{NFI}_m \end{bmatrix} \quad (13)$$

安全功能链 NFC 是对不同的安全功能 NF 进行编排, 所形成的具有特定顺序的安全功能的集合。因此, 可以认为安全功能链 NFC 是一种特定的安全管控形式。将安全控制结果 Ctrl 以安全功能链 NFC 的形式表征如下, 其中 NF_{ij} 表示在安全功能链中第 i 个位置, 且为安全功能池 NFP 中第 j 个编号的特定网络安全功能。

$$\text{Ctrl} \triangleq \text{NFC} \triangleq \{\text{NF}_{1,a}, \text{NF}_{2,b}, \dots, \text{NF}_{i,j}, \dots\} \quad (14)$$

在内生恶意行为检测框架中, 构造差异化的安全管控结果与用户行为检测结果密不可分。因此, 将安全控制结果表征为用户行为检测结果和安全功能池共同映射的形式, 如式 (15) 所示, 其中 Ctrl_i 、 NFC_i 分别表示第 i 个用户的安全管控



结果和生成的安全功能链。 $\sigma()$ 为安全功能链编排映射函数。

$$\begin{bmatrix} \text{Ctrl}_1 \\ \dots \\ \text{Ctrl}_i \\ \dots \\ \text{Ctrl}_n \end{bmatrix} \triangleq \begin{bmatrix} \text{NFC}_1 \\ \dots \\ \text{NFC}_i \\ \dots \\ \text{NFC}_n \end{bmatrix} \triangleq \sigma \left(\begin{bmatrix} \text{Char}_1^t \\ \dots \\ \text{Char}_i^t \\ \dots \\ \text{Char}_n^t \end{bmatrix}, \text{NFP} \right) \quad (15)$$

此外，安全控制模块还能够根据安全行为知识库的强大推理泛化能力，对接入用户进行快速决策。安全行为知识库 BKD 根据接入用户的用户身份标识 UID、用户历史行为集合 Hist_Behav、用户实时通信行为特征 RT_Char 等信息进行推理演化，能够实现安全控制结果的快速决策。将这一快速决策过程表征如式 (16)，其中 $\rho()$ 表示知识推理、映射决策函数。

$$\text{Ctrl} \triangleq \rho(\text{UID}, \text{Hist_Behav}, \text{RT_Char}, \text{BKD}) \quad (16)$$

4 自学习、自成长的恶意行为检测机制

在内生恶意行为检测框架中，通过行为检测模块和本地安全行为知识库的协同处理，实现了恶意行为的分布式检测与安全行为知识库动态更新，赋予了信息网络自学习、自成长的能力。

内生恶意行为检测框架所赋予的“自学习、自成长”能力主要体现在以下两个方面。首先，通过在跨域边界路由器中部署分布式恶意行为检测代理，构建基于分布式机器学习^[8]的恶意行为检测方法，通过本地训练、知识聚合、分布式感知学习、多域信息协同共享的方式，赋予行为检测模块自主学习、不断成长的能力。具体步骤如下。

步骤 1 部署在各个域的分布式恶意行为检测代理将本地收集到的用户实时通信行特征 RT_Char 作为本地分布式机器学习的训练数据。

步骤 2 恶意检测代理根据先验检测模型和本地训练数据 RT_Char 进行基于分布式机器学习

算法的恶意行为检测，得到该轮次的检测结果 Char 和经过本地训练后更新的参数。

步骤 3 各个域内的分布式检测代理将训练结束后的参数上传至知识聚合服务器，进行知识聚合，并更新恶意行为感知模型。

步骤 4 分布式检测代理对下发的模型和参数进行准确性验证后，更新本地模型和参数。

通过以上步骤，分布式行为检测模块能够自主地学习到其他域中的恶意行为信息，并通过知识聚合、参数共享的方式不断地更新本地检测模型。同时，行为检测模块充分利用分布式机器学习强大的泛化能力，具备检测未知恶意行为的可能，使其拥有不断成长的检测能力。

其次，“自学习、自成长”能力还体现在基于知识图谱^[9]技术构建的本地安全行为知识库中。本地安全行为知识库通过构建“用户身份-通信行为-用户信誉-安全控制”的链接关系，为恶意通信行为的检测和管控提供特征抽取、快速决策的能力，同时通过建立动态更新的域间信任链，赋予了本地行为知识库不断学习、成长的潜能。具体表现如下。

(1) 本地安全行为知识库根据从身份认证模块、信誉评估模块、行为检测模块以及安全控制模块中所抽取到的 UID、Hist_Behav、RT_Char、Ctrl 等特征信息，建立表征用户行为特征的行为知识 Behav_K，并基于知识图谱技术学习各特征之间的关系，构建具有连接关系的本地安全知识库 BKD。

(2) 本地安全知识库 BKD 利用知识图谱强大的知识推理能力，根据新的接入用户的特征信息基于知识图谱技术进行快速推演，做出快速决策，对用户通信行为进行管控。

(3) 此外，通过构建域间信任链，本地行为知识库能够定期与跨域行为知识库进行动态知识交换与更新，不断学习新的用户行为知识，更新本地安全行为支持库，使得本地行为知识库具备

自主学习、动态成长的特性。

5 原型系统搭建与实验

基于本文所提出的信息网络内生恶意行为检测框架，在现实环境中部署了基于该检测框架的原型系统并进行了初步的实验，对所提出的内生恶意行为检测框架的可行性进行验证。下面将围绕原型系统搭建以及实验验证两方面分别进行介绍。

图 2 为基于内生恶意行为检测框架的原型系统。在接入网关 A、B 处，部署了面向不同域的基于联邦学习算法的分布式恶意行为检测代理，以检测来自不同域的恶意通信行为。同时，在接入网关处部署了身份认证模块、信誉评估模块以及安全控制模块。此外，还分别为接入网关 A、B 部署了本地安全行为知识库 C、D，用于构建基于“用户身份-通信行为-用户信誉-安全控制”的链接关系。本地安全知识库 C、D 之间周期性进行信息更新，实现知识库动态成长。在核心网中部署了多种具有认证、聚合、流量清洗等特定功能的服务器。需要说明的是，知识聚合服务器针对不同的分布式检测代理所学习到的不同知识进行聚合、更新，赋予分布式检测代理自主学习的能力。

为了验证行为检测模块中分布式检测代理的性能，在接入网关 A、B 中的行为检测模块中部署了基于逻辑回归的联邦学习算法 FL，与此同时，在网关 A、B 处分别部署了相同的基于逻辑

回归的机器学习算法 LRA 和 LRB 作为对比算法。分别选用 CICDDoS 2019、UNSW-NB15 和 NSL-KDD 3 种恶意行为攻击数据集从两个不同的接入网向接入网关 A、B 发起模拟攻击，通过比较不同算法的检测结果来验证分布式检测代理的性能。图 3 比较了 3 个不同算法场景下的不同检测代理对恶意行为的检测结果。从图 3 中可以看出，对于 CICDDoS 2019、UNSW-NB15 和 NSL-KDD 3 种不同的恶意行为攻击数据集，基于联邦学习的分布式检测代理均表现了很高的恶意行为检测率，分别达到了 98.73%、81.34% 和 92.16%。此外，与其他两种只在单一接入网关处部署的机器学习算法 LRA 和 LRB 相比，基于联邦学习的分布式检测代理对恶意行为的检测更为精准。这是因为，基于联邦学习的分布式检测代理能够通过知识聚合、参数共享，综合地学习到所有接入域的恶意行为，而不是仅感知学习单个域内的恶意行为。

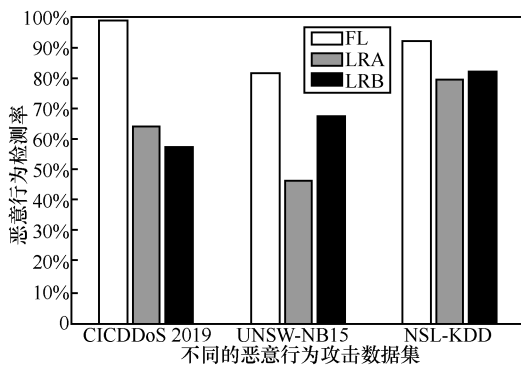


图 3 不同场景中行为检测模块对恶意行为的检测结果

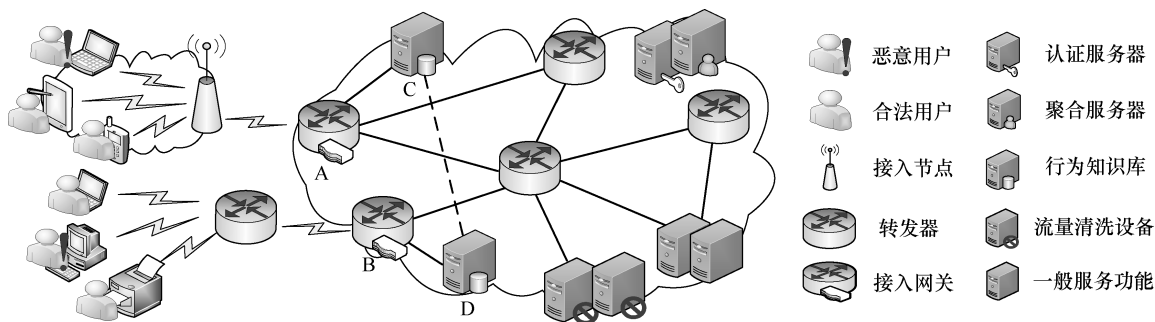


图 2 基于内生恶意行为检测框架的原型系统

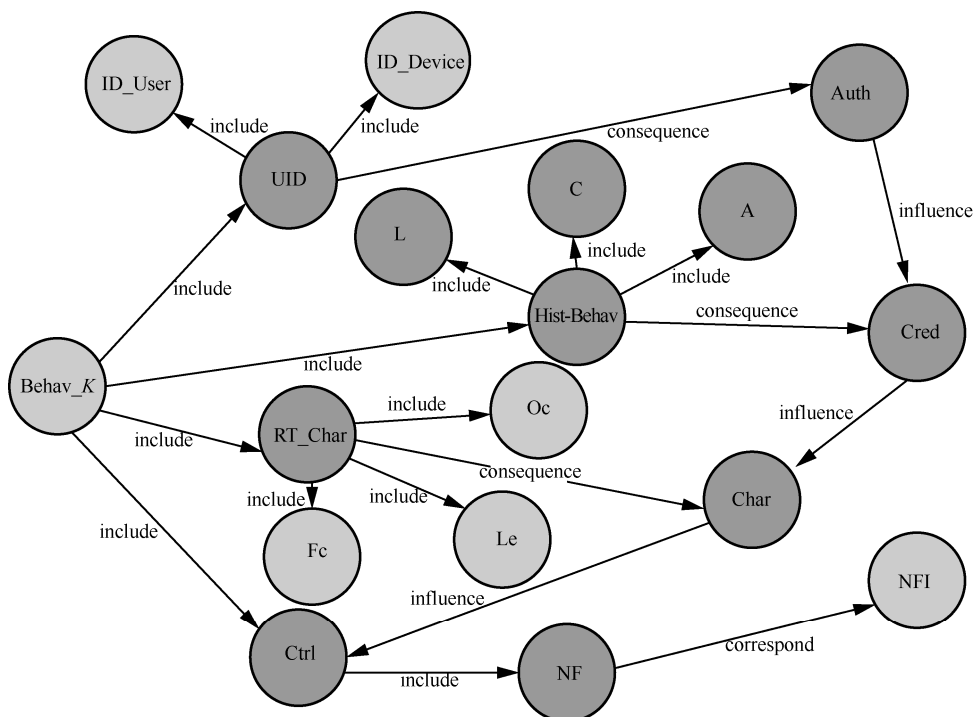


图4 表征用户行为链接关系的本地行为知识库

对所提框架中的安全行为知识库功能进行了验证, 图4为基于恶意行为特征构建的表征链接关系的本地行为知识库C。本地行为知识库C基于图数据库Neo4j, 根据收集部署在接入网关A中的身份认证、信誉评估、行为监测模块等模块中所抽取到的特征信息, 建立起了“用户身份-通信行为-用户信誉-安全控制”的链接关系。图4中“include”“consequence”“influence”和“correspond”分别表征实体之间的“包含”“结果”“影响”和“对应”的关系。可以根据图4所表征出来的不同用户恶意行为的链接关系, 对未出现或已出现过的恶意行为进行关系推演、快速决策等操作, 实现对用户通行行为进行快速管控。

6 结束语

为了从根本上阻断层出不穷的网络恶意行为, 构建安全可信的信息网络。本文从“内生安全”的角度, 提出了内生恶意行为检测框架, 对检测框架中的五大安全组件进行了建模分析, 介绍了自学习、

自成长的恶意行为检测机制。此外, 还对基于本文所提出框架的原型系统搭建方法与初步的实验结果进行了展示。希望通过本文的讨论, 推动信息网络内生恶意行为检测方法进一步的研究和探索。

参考文献:

- [1] Imperva's Threat Research Lab. Bad bot report 2020: bad bots strike back[R]. 2020.
- [2] 郭江兴. 新型网络技术发展思考[J]. 中国科学: 信息科学, 2018, 48(8):1102-1111.
WU J X. Thoughts on the development of novel network technology[J]. Scientia Sinica Informationis, 2018, 48(8): 1102-1111.
- [3] 刘杨, 彭木根. 6G 内生安全: 体系结构与关键技术[J]. 电信科学, 2020, 36(1): 11-20.
LIU Y, PENG M G. 6G endogenous security: architecture and key technologies[J]. Telecommunications Science, 2020, 36(1): 11-20.
- [4] 季新生, 黄开枝, 金梁, 等. 5G 安全技术研究综述[J]. 移动通信, 2019, 43(1): 34-39, 45.
JI X S, HUANG K Z, JIN L, et al. Overview on 5G security technology[J]. Mobile Communications, 2019, 43(1): 34-39, 45.
- [5] 江伟玉, 刘冰洋, 王闯. 内生安全网络架构[J]. 电信科学, 2019, 35(9): 20-28.
JIANG W Y, LIU B Y, WANG C. Network architecture with intrinsic security[J]. Telecommunications Science, 2019, 35(9): 20-28.
- [6] 唐燕群, 李为, 张立健, 等. 基于无线信道特征的内生安全通信技术及应用[J]. 无线电通信技术, 2020, 46(2): 159-167.

- TANG Y Q, LI W, ZHANG L J, et al. Endogenous secure communication based on characteristics of wireless channel[J]. Radio Communications Technology, 2019, 46(2): 159-167.
- [7] 张杰. 内生安全光通信技术及应用[J]. 无线电通信技术, 2019, 45(4): 337-342.
- ZHANG J. Technologies and applications of endogenously secure optical communication[J]. Radio Communications Technology, 2019, 45(4): 337-342.
- [8] NGUYEN T D, MARCHAL S, MIETTINEN M, et al. D³IoT: a federated self-learning anomaly detection system for IoT[C]//Proceedings of 2019 IEEE 39th International Conference on Distributed Computing Systems (ICDCS). Piscataway: IEEE Press, 2019: 756-767.
- [9] JIA Y, QI Y, SHANG H, et al. A practical approach to constructing a knowledge graph for cybersecurity[J]. Engineering, 2018, 4(1): 53-60.

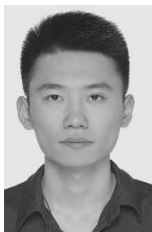
[作者简介]



涂哲(1994-), 男, 北京交通大学博士生, 主要研究方向为人工智能、空间网络、网络安全等。



周华春(1965-), 男, 博士, 北京交通大学教授、博士生导师, 主要研究方向为未来互联网体系架构、移动互联网、网络安全、空间网络等。



李坤(1997-), 男, 北京交通大学博士生, 主要研究方向为空间网络、人工智能、网络安全等。



王玮琳(1998-), 女, 北京交通大学硕士生, 主要研究方向为知识图谱、人工智能、网络安全。