

车载自组网中基于信任管理的安全组播协议设计

夏 辉^{1),2),5)} 张三顺^{1),2)} 孙运传³⁾ 肖 甫⁴⁾ 李 晔²⁾ 成秀珍⁵⁾

¹⁾(青岛大学计算机科学技术学院 山东 青岛 266000)

²⁾(山东省计算中心(国家超级计算济南中心) 山东省计算机网络重点实验室 济南 250014)

³⁾(北京师范大学经济与工商管理学院 北京 100875)

⁴⁾(南京邮电大学计算机学院 南京 210000)

⁵⁾(山东大学计算机科学与技术学院 山东 青岛 266237)

摘 要 路由协议是车载自组网(Vehicular Ad-hoc Networks, VANETs)在交通信息物理融合系统(Transportation Cyber Physical System, T-CPS)中有效运行的先决条件。作为 VANETs 移动信息和通信技术的基础,这些路由协议需要建立高效的通信路径以保障车辆间及时可靠的数据传输。然而车辆节点的快速移动性导致网络拓扑结构易变,增加了 VANETs 路由协议设计的难度。对计算、通信和控制技术的过分依赖,使得 VANETs 在缺乏公钥基础设施的情况下极易遭受各种类型的恶意攻击,信息传输的安全性受到极大威胁。为了解决路由安全问题,本文提出了一种抗攻击的轻量级可信安全组播路由协议(Trust-based Secure Multicast Routing Protocol, TSMRP)。首先,本文设计了一个高效的信任计算模型,模型综合直接信任和推荐信任两种信任决策因子得出车辆节点的总体信任值。在直接信任模型的构造过程中,利用基于波动类型识别的灰色马尔可夫信任预测算法准确得到车辆节点的直接信任值。在推荐信任模型的构造过程中,根据推荐车辆节点与监控者的交互关系,将推荐信任划分为三类,利用基于熵权的模糊层次分析法得出三种推荐信任的最优权重。其次,利用 NetLogo 仿真平台验证信任计算模型的有效性和准确性,其可以有效地提高恶意车辆节点的检出率和识别准确度。直接信任模型不仅可以应对常见的黑/灰洞攻击,而且可以有效地识别时序变换类攻击。基于反馈机制,推荐信任模型能够有效地应对抵毁攻击和共谋攻击。随后,本文设计实现了一种基于信任的安全组播协议,该协议可以在路由建立和路由维护过程中有效地识别和抑制恶意车辆节点,并通过排除恶意中继节点进而建立安全可靠的通信路径。此外,为了进一步提高路由效率,在路由维护阶段本文提出了两种改进机制。转发组节点复用机制和边缘路径删除机制能够有效地减少转发组节点的个数和路由请求报文的洪泛,从而减少了信任模型融合带来的路由开销,避免了链路阻塞造成的大量数据包丢失,进而使得新协议更加适应高速移动环境。最终,详实的实验结果表明,相较于相关协议,TSMRP 提高了数据投递率,降低了路由开销、平均端到端延迟和单位信包量。

关键词 车载自组网;安全路由协议;恶意攻击;信任计算;灰色马尔可夫预测

中图法分类号 TP393

DOI号 10.11897/SP.J.1016.2019.00961

Design of Trust-Based Secure Multicast Routing Protocol in VANETs

XIA Hui^{1),2),5)} ZHANG San-Shun^{1),2)} SUN Yun-Chuan³⁾ XIAO Fu⁴⁾ LI Ye²⁾ CHENG Xiu-Zhen⁵⁾

¹⁾(College of Computer Science and Technology, Qingdao University, Qingdao, Shandong 266000)

²⁾(Shandong Computer Science Center (National Supercomputer Center in Jinan),
Shandong Provincial Key Laboratory of Computer Networks, Jinan 250014)

³⁾(Business School, Beijing Normal University, Beijing 100875)

⁴⁾(College of Computer, Nanjing University of Posts and Telecommunications, Nanjing 210000)

⁵⁾(Department of Computer Science and Technology, Shandong University, Qingdao, Shandong 266237)

Abstract The routing protocol is a prerequisite for the effective operations of vehicular ad hoc

收稿日期:2018-06-19;在线出版日期:2019-03-18。本课题得到国家自然科学基金(61872205,61371185)、国家自然科学基金重点项目(61832012)、国家重点研发计划(YFB0803400)、国家千人计划、山东省高等学校科技计划项目(J16LN06)、青岛市应用基础研究计划(18-2-2-56-jch)、国家留学基金和山东省计算机网络重点实验室开放课题基金(SDKLCN-2018-07)资助。夏 辉,博士,副教授,中国计算机学会(CCF)会员,主要研究方向为物联网、无线网络、信息安全、可信计算。E-mail: xiahui@qdu.edu.cn。张三顺,硕士研究生,主要研究方向为物联网、无线网络、信息安全。孙运传,博士,教授,中国计算机学会(CCF)高级会员,主要研究领域为数据科学、物联网、信息安全。肖 甫,博士,教授,中国计算机学会(CCF)高级会员,主要研究领域为物联网与传感网、网络通信协议与算法。李 晔,博士,研究员,中国计算机学会(CCF)高级会员,主要研究领域为多媒体信号处理、无线通信。成秀珍,博士,教授,中国计算机学会(CCF)高级会员,主要研究领域为无线网络与通信、物联网、隐私保护、分布式算法。

networks (i. e. , VANETs) in the transportation cyber-physical system (i. e. , T-CPS). As the basis of mobile information and communication technologies in VANETs, these protocols must be able to establish effective and efficient routing paths to guarantee timely and reliable data transmission between vehicles. However, the rapid movement of vehicles leads to frequent changes in topology and increases the difficulty of designing routing protocols. Due to the excessive reliance on computing, communication and control technologies, VANETs are vulnerable to various types of malicious attacks in the absence of public key infrastructure. And the security of information transmission is significantly threatened. To tackle the routing security issue, this paper proposes an attack-resistant light-weight trust-based secure multicast routing protocol (i. e. , TSMRP), which can resist multiple malicious attacks based on the integration of trust. This paper first designs a highly efficient trust computing model, which integrates two trust decision factors, as the direct trust and the recommendation trust, to derive the overall trust value of a vehicle. During the construction of direct trust model, the grey Markov prediction algorithm based on the fluctuation type recognition is utilized to calculate the direct trust value of a specific vehicle accurately. Meanwhile, during the construction of recommendation trust model, the recommending vehicles are classified into three types based on their interaction relationship with the monitoring vehicle. And further, the fuzzy analytic hierarchy process based on entropy mechanism is used to assign optimal weights for these three types of recommendation trust. Secondly, the NetLogo simulation platform is used to verify the effectiveness and accuracy of this trust computing model, which can effectively improve the detection rate and the identification accuracy of malicious vehicles. This direct trust model can not only deal with the common black/grey hole attacks but can also effectively identify the timing-shift attacks. The recommendation trust model can effectively deal with slander attacks and collusion attacks. Thirdly, a trust-based secure multicast protocol (i. e. , TSMRP) is designed to effectively identify and suppress malicious vehicles during the route establishment phase and the route maintenance phase. Thereby, this new trust-based routing protocol can develop communication paths of high security and reliability via excluding the malicious relay vehicles. Besides, to further improve routing efficiency, this paper proposes two improvement mechanisms including the forwarding group node reuse mechanism and the edge path deletion mechanism during the route maintenance phase. These two mechanisms can reduce the number of forwarding group nodes and the flooding of routing request packets effectively, thus reduce the routing overhead caused by trust model fusion and avoid the loss of a large number of packets due to the link blocking. These beneficial effects make this new proposed protocol more adaptable to the environment with high mobility. Finally, compared with the other relevant protocols, the detailed experimental results show that the TSMRP improves the packet delivery ratio. On the contrary, the routing overhead, the average end-to-end latency and the byte sent per byte delivered are reduced.

Keywords vehicular ad hoc network; secure routing protocol; malicious attacks; trust computing; grey Markov prediction

1 引 言

随着汽车制造商之间的竞争日益激烈,车载自组网(Vehicular Ad Hoc Networks, VANETs)正经

历着巨大的演变. 导航系统、嵌入式传感器和新标准化通信技术的引入,极大地改善了人们的生活水平. 汽车在行驶过程中可以通过交换驾驶员行驶意图的信息来避免碰撞,应急车辆可以提前警示前方的车辆做出避让操作^[1]. 然而车辆的高速移动性会导致

基于车载网基础设施所建立的链路频繁切换,甚至造成断裂。在道路边缘部署多个固定信号收发装置的成本过高,除去数据传输的自身成本外,部署车载蜂窝网(Vehicular Cellular Network, VCN)和车载无线局域网(V-WLAN)还需要使用昂贵的通信收发设备^[2]。因此, VANETs 内部通信过程主要采用车辆到车辆(Vehicle-to-Vehicle, V2V)的模式,不需要配套基础设施的参与。车辆依靠无线收发器与其它相邻车辆交换数据,然后将数据分组经由相邻车辆路由转发到不在通信范围内的目的车辆。

数据分发是 VANETs 领域的一个重要研究方向,其主要依赖于高效的轻量级路由协议来实现^[3]。为满足不同用户的需求, VANETs 路由应具有一定的约束条件,例如较低的端到端延迟和丢包率。然而车辆的高速移动性会导致所选择的初始路由路径迅速改变,引起数据传输延迟增大和数据包丢失,因此 VANETs 路由协议必须具有较高的鲁棒性,能够快速适应网络拓扑的频繁变化。此外,证书颁发机构(Certificate Authority, CA)等公钥基础设施的缺失,令 VANETs 更易遭受各种类型的恶意攻击和安全威胁。安全路由的建立,能够保证数据在车辆节点之间可信传输,使司机能够根据可靠的信息做出正确的决策。为了实现路由安全,传统方法是采用基于密码学的安全机制^[4-6],但其存在明显缺陷:(1)可信中央控制节点的缺失使得 VANETs 系统节点间的信任和授权较难实现;(2)PKI 基于集中式服务器,专注于路由数据包的完整性检测以防止恶意实体对非可变字段的修改;(3)此类机制无法识别网内拥有合法授权身份的恶意实体,易遭受 DoS/DDoS 攻击、女巫攻击和假冒攻击等内部攻击;(4)安全协议执行过程中对数据包或字段的频繁加解密操作将消耗大量的带宽,产生高昂的网络开销。相较而言,基于信任管理的安全路由机制被认为是一种更有前途的解决方案,能有效弥补前者存在的缺陷。该类机制能够隔离网内行为不端的恶意节点,保证车辆之间的安全交互对上下文情境感知、数据传输和融合的可靠性和有效性起到关键作用。

车载自组网可以被看作一类较为特殊的移动自组网(Mobile Ad Hoc Networks, MANETs),二者具有许多共性^[7]。MANETs 中基于网格的路由协议(如 ODMRP, NSMP, OBAMP 等)具有较高的鲁棒性,能够有效适应 VANETs 中车辆的快速移动性。因此,本文借鉴网格结构的优势,提出了一种以信任计算模型为基础的轻量级安全组播路由协议

(Trust-based Secure Multicast Routing Protocol, TSMRP)。本文主要贡献如下:

(1)考虑到车辆节点间历史交互数据的波动性,提出了基于波动类型识别的灰色马尔可夫预测模型,对直接信任进行预测。该模型不仅可以应对普通的黑/灰洞攻击,而且利用波动趋势的概念能有效识别节点发起的时序变换类攻击;

(2)根据交互关系将推荐节点进行系统地分类,利用基于熵权的模糊层次分析法为各类推荐节点的推荐信息分配最佳权重。其中,基于反馈机制提出了一种推荐可信度的计算方法,能有效地应对诋毁攻击和共谋攻击,过滤不良推荐信息;

(3)设计并实现了适用于 VANETs 的轻量级安全路由协议 TSMRP。本文将信任融入组播路由协议的设计过程中,使其能够抵御多种恶意攻击;

(4)为了进一步提高路由效率,在路由维护阶段提出了两个有效的机制,分别是转发组节点复用机制和边缘路径删除机制,它们能够有效地减少转发组节点的个数,减少路由请求报文的洪泛,从而降低路由开销,避免因链路阻塞而导致数据包大量丢失。

本文第 2 节介绍目前 VANETs 路由协议和信任管理的相关工作;第 3 节详细地描述信任计算模型;第 4 节对信任模型的性能进行仿真验证;第 5 节详细介绍基于信任管理的轻量级安全组播路由协议 TSMRP;第 6 节对新协议进行仿真模拟和性能对比;第 7 节对全文进行总结,提出下一步研究计划 and 方向。

2 相关工作

2.1 VANETs 路由协议

路由协议是 VANETs 运行的基础,近年来研究者们为设计出高效的路由协议做了大量研究。Togou 等人^[8]提出了分布式路由协议 SCRP,即网络中的节点根据收集到的延迟和连接信息为路径分配权重,选择具有最低聚合权重的路径来转发数据分组。为了解决个体车辆到路边单元端到端传输性能不可靠的问题, Wang 等人^[9]提出了一种新的分组转发方案,即利用车辆的移动性差异来处理最佳网络吞吐量,以平衡网络中的数据流量。Lin 等人^[10]设计了一个基于移动区域的架构,即车辆彼此协作形成动态移动区域,以促进信息传播。随后又将移动对象建模和索引技术从大型移动对象数据库引入到 VANETs 路由协议的设计中,提高了协议的效率。

Wu 等人^[11]考虑了数据的传输速率和多跳传输效率对 VANETs 路由协议性能产生的影响,提出一种能够通过与环境交互确定最佳传输参数的路由协议.文献[12]中,作者提出了微观拓扑学(MT)的新概念,设计了一种基于 MT 的新型街道中心路由协议 SRPMT. Zhu 等人^[13]研究发现车辆在多级场景中移动时呈现出多级特征,从而提出了一种面向多级场景的贪婪机会路由协议 M-GOR,该协议利用连通概率的计算方法和贪婪机会转发算法来应对多层结构的影响. Yao 等人^[14]通过车辆运动表现出的高度重复性和历史运行路线,利用隐马尔可夫模型预测车辆的未来位置,提出了一种新颖的预测路由协议 PRHMM.

上述路由协议的设计大多基于节点友好合作的网络环境,侧重于维护车辆节点之间的正常通信,而很少考虑恶意车辆对协议性能的影响.为了有效应对潜在的多种恶意攻击,识别恶意车辆并抑制它们对网络造成的负面影响,需要设计实现一种高效、可靠且轻量级的安全路由协议.

2.2 VANETs 信任管理

与基于密码体系的硬安全机制相比,基于信任管理的软安全机制通过另外一条途径保证了系统的安全性. Kerrache 等人^[15]在文中将 VANETs 中路由的安全解决方案分为两类:基于密码学的安全解决方案和基于信任管理的安全解决方案.文中对密码学方法和信任管理方法进行了对比,阐明了信任管理方法的多种优势.文献[16]提出了一种基于多 QoS 约束的可靠感知路由算法,该算法提供了一种切实可行的方法来抵御路由攻击.在算法中,采用蚁群优化技术,根据数据流量类型决定多个 QoS 约束条件来选择可行路由,并利用面向 VANETs 的演化图模型对路由控制消息进行检测.在文献[17]中 Li 等人提出了一种针对 VANETs 的抗攻击信任管理方案 ART,该方案不仅能够检测和处理恶意攻击,而且可以评估数据和移动节点的可信性.其中数据信任利用多个车辆感测和收集的数据来评估,节点信任利用功能信任度和推荐信任度进行评估. Dietzel 等人^[18]提出可以通过多跳路由协议中的冗余信息传播来检查数据一致性.并介绍了三种基于图论的度量标准,衡量协议的冗余度,并将度量指标应用到路由协议中,从而检测具有攻击行为的车辆,并将它们排除出网络. Krishna 等人^[19]对已有的车辆不正当行为检测算法进行改进,通过共识机制来计算协作节点的可信度,提出了 k 均值聚类算法来过滤不可

信推荐信息以提高信任计算的精确度. Marmol 等人^[20]提出了基于信任和声誉的安全方案,快速准确地区分恶意或自私节点在网络中传播的虚假信息,避免因接收到虚假信息而做出错误决策的举动.

然而上述安全路由机制中信任模型的构建均存在一定的局限性,不能有效地抵御多种恶意攻击.基于信任管理的安全路由机制仍旧是当前研究的一个热点方向.

3 信任模型

基于信任管理的安全路由机制的基础是高效的信任计算模型,依靠信任模型收集信任评估信息,快速准确地识别并抑制网内恶意节点,以有效应对各类攻击.本节所提出的信任计算模型全面地考虑了直接信任与推荐信任这两种核心决策因子.

3.1 直接信任模型

传统的信任计算模型在直接信任因子的计算过程中,往往采用贝叶斯算法、神经网络、机器学习或 β 分布等方法对直接信任值进行预测.然而当预测样本较小、波动幅度较大时,预测的精确度会大幅降低.为此,夏怒等人^[21]提出一种基于波动类型识别的行为预测算法 FTI-RNBP,其以灰色模型为基础,能够利用少量的样本数据预测节点的行为值.囿于灰色模型的自身局限性,该算法预测精度不高,而且其不能有效识别恶意节点发起的时序变换类攻击(如 on-off 攻击).本文对该算法进行大幅改进,提出了基于波动类型识别的灰色马尔可夫预测算法 FTI-MSCGM,新算法弥补了前者的不足,能够更加精确地预测节点的直接信任值.

3.1.1 波动类型识别

为方便描述,本文将两个节点(i 和 j)的交互历史均分为 t 个周期.在第 k 个周期内,将节点 i 监测到节点 j 的转发率定义为节点 j 的行为值 $A_{i,j}(k)$.假设 n 为第 k 周期内节点 j 与节点 i 交互过程中收到的数据包个数, m 为节点 j 实际转发的数据包个数,则行为值计算如下:

$$A_{i,j}(k) = m/n \quad (1)$$

然而如果某周期内节点间交互次数过少,则通过转发率得到的行为值误差偏大.在行为值计算过程中应考虑当前周期内节点间的交互次数,交互次数越多,通过转发率得到的行为值误差越小.本文使用活跃度函数 $\varphi(x)$ 表示交互次数对节点行为值的影响程度.该函数随交互次数的增加单调递增,并趋

近于 1. 计算公式如下：

$$\varphi(x) = \frac{\arctan(x-a)}{\pi} + 0.5 \quad (2)$$

其中, x 为节点间的交互次数, 参数 a 为交互次数调节因子, 合理的取值应为介于区间 $[3, 7]$ 的整数. 利用活跃度函数 $\varphi(x)$ 对节点 j 的行为值进行调整, 计算如下：

$$A_{i,j}(k) = \begin{cases} \varphi(n) \frac{m}{n} + (1 - \varphi(n)) A_0, & n \neq 0 \\ A_0, & n = 0 \end{cases} \quad (3)$$

其中, A_0 为节点的缺省行为值, 表示节点间没有交互时的初始行为值. 由此可以得到历史交互过程中各个评估周期内节点 j 的行为值 $A_{i,j}(1), A_{i,j}(2), \dots, A_{i,j}(t)$. 依靠历史交互信息计算节点 i 对节点 j 的直接信任, 即根据节点 j 的 t 个历史行为值来预测 j 在第 $t+1$ 周期的行为值 $A_{i,j}(t+1)$ ：

$$TD_{i,j} = A_{i,j}(t+1) \quad (4)$$

定义 1. 级比序列. 节点 j 的所有行为值构成了序列 $A = \{A_{i,j}(1), A_{i,j}(2), \dots, A_{i,j}(t)\}$, 序列 A 中相邻行为值的比值构成了级比序列 $\gamma = \{\gamma_1, \gamma_2, \dots, \gamma_{t-1}\}$, 其中 $\gamma_l = A_{i,j}(l+1)/A_{i,j}(l)$.

依据灰色模型若行为序列 A 的级比序列 γ 满足 $|\max(\gamma_l) - \min(\gamma_l)| \leq \theta$, 则称此序列 A 为平滑序列. 其中 θ 为常数, 级比序列 γ 动态性取值：

$$\theta = \frac{1}{t-1} \sum_{l=1}^{t-1} |\gamma_l - \min(\gamma)| \quad (5)$$

根据限制条件 $|\max(\gamma_l) - \min(\gamma_l)| \leq \theta$, 对序列 A 中的行为值进行筛选. 本文将筛选出的异常值称为波动值, 进而组成波动序列 F . 可将波动值分为以下两类：

(1) 突变波动. 仅干扰当前时刻的行为值, 对后续的行为值序列无影响；

(2) 迁移波动. 不仅干扰当前时刻的行为值, 而且对后续的行为值序列有影响. 如波动值 $A_{i,j}(2)$, 若与它相邻的下一个波动值为 $A_{i,j}(5)$, 则认为 $A_{i,j}(2)$ 的影响力大于两个数据, 称为迁移波动；否则若下一个相邻波动值为 $A_{i,j}(3)$, 则认为 $A_{i,j}(2)$ 的影响力小于两个数据, 称为突变波动. 然而仅凭借单一波动值无法直接对其波动类型进行识别, 本文借鉴灰色马尔可夫时序预测模型^[22]来分析历史波动值的变化规律, 从而预测当前波动所属类型.

3.1.2 预测算法思想

本节提出波动趋势的概念, 可有效应对 on-off 攻击等时序变换类攻击. 对于每一个识别出的波动值, 都对应一个波动趋势 $\mu(t)$: 若该波动值大于前

一个行为值, $\mu(t)$ 值置为 1; 若该波动值小于前一个行为值, $\mu(t)$ 值置为 0. 序列 μ 为波动趋势构成的序列, 根据波动值类型及波动趋势序列 μ 的不同, 提出以下两种预测方法：

(1) 若识别出的波动值 $A_{i,j}(t)$ 全部为迁移波动 (波动值个数不小于 2) 且序列 μ 中 0, 1 交替出现, 则认为节点较大概率发起时序变换类攻击 (如 on-off 攻击). 在这种情况下, 若直接根据历史行为值进行信任预测, 可能会做出错误的判断. 因此要将节点表现良好时的行为值排除, 利用节点表现较差时的行为值集合 $A^* (A^* \subset A)$ 进行信任值的预测. 计算方法如下：

$$A_{i,j}(t+1) = \text{MGPred}(A^*) \quad (6)$$

函数 $\text{MGPred}()$ 表示灰色马尔可夫模型预测过程.

(2) 若 $A_{i,j}(t)$ 为突变波动, 表示 $A_{i,j}(t)$ 和节点 j 的下一个周期的行为值 $A_{i,j}(t+1)$ 之间存在较大差异. 但 $A_{i,j}(t+1)$ 与序列中的其它非波动值相似, 因此可以使用灰色马尔可夫模型直接进行计算：

$$A_{i,j}(t+1) = \text{MGPred}(A) \quad (7)$$

若 $A_{i,j}(t)$ 为迁移波动或者不是波动值, 则表明 $A_{i,j}(t+1)$ 与 $A_{i,j}(t)$ 属于同一个平滑区间. 可使用灰色马尔可夫模型对平滑级比序列的序列值进行预估, 再结合当前行为值 $A_{i,j}(t)$ 对 $A_{i,j}(t+1)$ 进行预测：

$$A_{i,j}(t+1) = A_{i,j}(t) \times \text{MGPred}(\gamma') \quad (8)$$

其中, γ' 是平滑级比序列, 即除去级比序列 γ 中的波动值.

3.1.3 灰色马尔可夫预测模型

灰色马尔可夫预测模型 Markov-SCGM(1,1) 基于系统云灰色预测模型 SCGM(1,1), 利用马尔可夫链的优势, 来确定 SCGM(1,1) 模型拟合精度指标时序的状态间转移规律^[22]. 该方法可修正灰色模型, 提高随机波动性较大的序列预测精度.

(1) SCGM(1,1) 模型建立

假设粗糙样本序列为 $X^{(0)} = \{x^{(0)}(1), x^{(0)}(2), \dots, x^{(0)}(n)\}$, 对 $X^{(0)}$ 进行积分变换, 得到序列 $X^{(1)} = \{x^{(1)}(1), x^{(1)}(2), \dots, x^{(1)}(n)\}$, 其中 $x^{(1)}(k) = \sum_{m=2}^k \bar{x}^{(0)}(m)$, $k = 2, 3, \dots, n$; $\bar{x}^{(0)}(k) = (x^{(0)}(k) + x^{(0)}(k-1))/2$.

假定序列 $X^{(1)}$ 与非齐次指数离散函数 $f_r(k) = be^{a(k-1)} - c$ 满意趋势关联, 则 SCGM(1,1) 模型表示为

$$\frac{d\hat{x}^{(1)}(k+1)}{dk} = a\hat{x}^{(1)}(k+1) + U, \quad k \geq 2 \quad (9)$$

其一次响应函数为

$$\hat{x}^{(1)}(k) = (\hat{x}^{(1)}(1) + U/a)e^{ak} - U/a,$$

其中, $\hat{x}^{(1)}(1) = b - c$, $U = ac$, a, b, c 的值如下:

$$a = \ln \frac{\sum_{k=3}^n \bar{x}^{(0)}(k-1) \bar{x}^{(0)}(k)}{\sum_{k=3}^n (\bar{x}^{(0)}(k-1))^2},$$

$$b = \frac{(n-1) \sum_{k=2}^n e^{a(k-1)} x^{(1)}(k) - \sum_{k=2}^n e^{a(k-1)} \sum_{k=2}^n x^{(1)}(k)}{(n-1) \sum_{k=2}^n e^{2a(k-1)} - \left(\sum_{k=2}^n e^{a(k-1)} \right)^2},$$

$$c = \frac{1}{(n-1)} \left[\left(\sum_{k=2}^n e^{a(k-1)} \right) b - \sum_{k=2}^n x^{(1)}(k) \right].$$

对 $\hat{x}^{(1)}(k)$ 进行还原处理, 得原始数据的 SCGM (1,1) 预测模型为

$$\hat{x}^{(0)}(k) = e^{a(k-1)} \frac{2b(1-e^{-a})}{1+e^{-a}} \quad (10)$$

令 $y(k) = x^{(0)}(k) / \hat{x}^{(0)}(k)$ 为灰拟合精度指标, 反映拟合值对原始数据序列的偏离程度.

(2) 基于模糊聚类的状态划分

目标节点可信度的变化趋势是一个非平稳随机过程, 不同的时间序列值都有其各自特殊的边界. 本文引入模糊聚类的方法来进行原始时间时序值的状态划分. 首先建立一个基本的聚类中心 c_i ; 基于这个中心, 可以设计几个并行区域来代表不同的状态集合. 基于以下等式执行相似度计算:

$$\mu_i[y(k)] = (1/\|y(k) - c_j\|^2) / \sum_{\lambda=1}^i (1/\|y(k) - c_\lambda\|^2) \quad (11)$$

然后, 基于上述步骤, 可以对所有获得的聚类中心进行迭代操作:

$$c_j = \sum_{j=1}^n (\mu_i[y(k)]^2 y(k)) / \sum_{j=1}^n \mu_i[y(k)]^2 \quad (12)$$

根据更新的聚类中心可以将每个原始样本重新分类为不同的聚类(状态), 并获得每个群集的边界 $[b_{i \min}, b_{i \max}]$.

(3) 自相关系数的计算

为了正确反映各阶对马尔可夫链预测值的影响, 采用 $y(k)$ 各阶自相关系数反映权值大小:

$$r_k = \frac{\sum_{l=1}^{n-k} (y(l) - \bar{y})(y(l+k) - \bar{y})}{\sum_{l=1}^n (y(l) - \bar{y})^2} \quad (13)$$

其中 $y(l)$ 代表第 l 次统计的灰色拟合水平, $\bar{y}$ 表示灰色拟合序列的平均值, n 代表序列的长度, k 表示马尔可夫链的阶. 对各阶自相关系数作标准化处理:

$$\theta_k = |r_k| / \sum_{k=1}^m r_k \quad (14)$$

所获得的值可以用作各阶马尔可夫链的权重, 参数 m 表示最大阶数.

(4) 马尔可夫链预测

根据状态的划分, 建立 k 阶 $P_{ij}^{(k)}$ 转移概率矩阵:

$$P^{(k)} = \begin{bmatrix} P_{11}^{(k)} & P_{12}^{(k)} & \cdots & P_{1n}^{(k)} \\ P_{21}^{(k)} & P_{22}^{(k)} & \cdots & P_{2n}^{(k)} \\ \vdots & \vdots & \vdots & \vdots \\ P_{n1}^{(k)} & P_{n2}^{(k)} & \cdots & P_{nn}^{(k)} \end{bmatrix} \quad (15)$$

依据上述矩阵, 经过 k 步, 状态 E_i 到达状态 E_j 的转移概率为

$$P_{ij}^{(k)} = m_{ij}^{(k)} / M_i \quad (16)$$

$m_{ij}^{(k)}$ 表示经过 k 步, 状态 E_i 转移到状态 E_j 的次数, M_i 表示状态 E_i 出现的次数. 需要注意在计算 M_i 的过程中, 要删除最后 k 个状态.

使用相应的转移概率矩阵, 可以发现初始状态的出现. 进而计算每行加权转换概率的总和:

$$P_i = \sum_{k=1}^m \theta_k P_i^{(k)} \quad (17)$$

最后, 根据 P_i 的最大值来确定未来状态. 当最终状态确定后, 使用线性插值法来计算灰色拟合精度指数, 公式如下:

$$\hat{y}(n+1) = b_{i \min} \times \frac{P_{i-1}}{P_{i-1} + P_{i+1}} + b_{i \max} \times \frac{P_{i+1}}{P_{i-1} + P_{i+1}} \quad (18)$$

(5) 预测值的确定

使用函数 $MGPred(A)$ 表示上述过程, 预测值计算如下:

$$MGPred(X^{(0)}) = \hat{x}^{(0)}(n+1) \times \hat{y}^{(0)}(n+1) \quad (19)$$

3.1.4 直接信任模型算法描述

该算法的时间复杂度等同于灰色马尔可夫预测模型的算法复杂度. 灰色马尔可夫预测模型的算法复杂度为 $O(n^3)$, 其中 n 为序列的长度, 而评估周期个数决定了序列的长度, 因此该算法的时间复杂度主要取决于划分的评估周期数 t . 该算法的网络开销主要来源于信任信息的收集过程, 而这些信息(如接收数据包量、转发数据包量和交易次数等)很容易在节点交互过程中获得, 因此, 直接信任模型产生的网络开销很小.

算法 1. 直接信任模型算法描述.

输入: $A = \{A_{i,j}(1), A_{i,j}(2), \dots, A_{i,j}(t)\}$

输出: $TD_{i,j}$

1. FOR $l \leftarrow 1$ to $t-1$ DO

2. $\gamma_l = A_{i,j}(l+1) / A_{i,j}(l)$;

```

3. END FOR
4. WHILE  $\gamma \neq \emptyset$  DO
5.   IF  $|\max(\gamma_i) - \min(\gamma_i)| \leq \theta$  THEN
6.     add  $\gamma_i$  to  $\gamma'$ ;
7.   ELSE
8.     add  $\gamma_i$  to  $F$ ;
9.   END WHILE
10. obtain  $\mu(t)$  corresponding to each  $\gamma_i$  in  $F$ ;
11. IF  $A_{i,j}(t)$  is migration fluctuation and 0, 1 alternate in  $\mu(t)$  THEN
12.   get  $A^*$  by excluding good behavior values in  $A$ ;
13.    $A_{i,j}(t+1) = \text{MGPred}(A^*)$ ;
14. ELSE IF  $A_{i,j}(t)$  is abrupt fluctuation THEN
15.    $A_{i,j}(t+1) = \text{MGPred}(A)$ ;
16. ELSE
17.    $A_{i,j}(t+1) = \text{MGPred}(\gamma')$ ;
18. END IF
19.  $TD_{i,j} = A_{i,j}(t+1)$ ;
20. RETURN  $TD_{i,j}$ ;

```

3.2 推荐信任模型

如果节点无法判断其感兴趣的节点是否可信(两者之间无交互或交互较少),该节点需要收集来自多个推荐节点的用于信任计算的推荐信息。但是对于不同类型的推荐节点,评估节点对其推荐信息可信度的判断标准也不同,如果采用相同的方法进行计算,会带来较大的误差。因此,本文基于推荐节点 k 与评估节点 i 的关联程度,将其分为三类并分别计算各自的推荐信任。最后通过基于熵权的层次分析法^[23]确定各类推荐信息的最优权重,得出总推荐信任值 $TR_{i,j}$ 。

3.2.1 推荐可信度

推荐节点提供的推荐信息是否可信直接影响到总推荐信任计算的准确性,本文使用推荐可信度 $C_{i,k}$ 表示评估节点 i 对推荐节点 k 给出的推荐信任的信任程度。传统的推荐可信度计算方法是利用节点 i 与节点 k 之间的交互关系(如直接交互历史、公共节点等)来计算,本文将这类基于交互关系得到的推荐可信度称为关系可信度 $RC_{i,k}$ 。该类可信度虽然在一定程度上能够衡量推荐信任的可靠程度,但是并不能有效地应对节点的诋毁攻击或共谋攻击,例如节点 k 在数据转发过程中一直是可信的,但总提供虚假的推荐信息,甚至与其它节点共谋提供不良推荐信息。由于上述涉及到的节点间的交互过程都是可信的,因此单纯地利用交互关系就会得到较高的关系可信度,无法识别这些不良推荐。为了解决上述问题,本文提出了基于反馈机制的推荐可信度计算方法,称为反馈可信度 $FC_{i,k}$ 。

当节点 i 和节点 j 完成交互后,根据交互的结

果,可以得到本次交互的真实信任值 $RT_{i,j}$ 。根据真实信任值来检验交互之前推荐节点 k 提供的推荐信任的准确度。本文定义一个偏离阈值 d , d 的合理取值范围应为 $0.1 \sim 0.2$ 。若 $|RT_{i,j} - TD_{k,j}| < d$,二者偏离较小,表明推荐节点 k 提供了正确推荐;若 $|RT_{i,j} - TD_{k,j}| \geq d$,二者偏离较大,表明推荐节点 k 提供了不良推荐。这种反馈推荐机制,可以较为准确地识别出不良推荐,避免伪装节点(如诋毁节点、共谋节点)通过与评估节点诚信交互获得较高的推荐可信度,随后不停地推荐不良信息。

本文定义反馈可信度取值 $FC_{i,k} \in [0, 1]$,利用反馈机制进行判断,具体过程为:若提供正确推荐,反馈可信度逐渐递增,但不超过 1;若提供不良推荐,反馈可信度逐渐递减,但不低于 0;若无反馈信息,反馈可信度保持不变 $FC_{i,k}$ 。反馈可信度计算公式如下:

$$FC_{i,k} = \begin{cases} FC_{i,k} + r, & |RT_{i,j} - TD_{k,j}| < d \\ FC_{i,k} - p, & |RT_{i,j} - TD_{k,j}| \geq d \\ FC_{i,k}, & \text{无推荐} \end{cases} \quad (20)$$

其中, r 为奖励因子, p 为惩罚因子;为了体现对不良推荐的惩罚,反馈可信度减小的速度大于增加的速度,即 $p > r$ 。

本文利用关系可信度和反馈可信度,综合得到节点的推荐可信度,计算如下:

$$C_{i,k} = \frac{RC_{i,k} + FC_{i,k}}{2} \quad (21)$$

其中 $RC_{i,k}$ 的取值范围为 $[0, 1]$,并根据推荐节点类型的不同分别进行计算。

3.2.2 直接推荐信任

本文将与评估节点 i 有过直接交互关系的推荐节点 k 定义为直接推荐节点,评估节点对该类推荐节点的直接信任值 $TD_{i,k}$ 可以直接作为该类节点的关系可信度 $RC_{i,k}$ 。利用式(21)得到推荐可信度 $C_{i,k}$ 。若 $C_{i,k} < 0.5$,评估节点将直接排除推荐节点提供的信任推荐信息。不同推荐者提供的直接推荐信息所产生的影响与各自的历史交互量成正向相关性,交互越多越频繁,则影响越大。综合考虑可信度和交互次数,直接推荐信任 $TR_{i,j}^d$ 计算如下:

$$TR_{i,j}^d = \sum_{k=1}^{N^d} \frac{C_{i,k} M_{i,k}}{\sum_{k=1}^{N^d} C_{i,k} M_{i,k}} \times TD_{k,j} \quad (22)$$

其中, N^d 为直接推荐节点的个数, $M_{i,k}$ 为节点 i 和直接推荐节点 k 的交互次数。

3.2.3 间接推荐信任

本文将与评估节点 i 有间接交互关系的推荐节

点 k 定义为间接推荐节点(譬如推荐节点 k 与评估节点 i 有重叠的交互节点). 可建模一个相似度函数 $Sim_{i,k}$ 来描述两者对公共问题的看法差异, 本文使用余弦相似度函数来实现, 计算如下:

$$Sim_{i,k} = \frac{\sum_s TD_{i,s} TD_{k,s}}{\sqrt{\sum_s TD_{i,s}^2} \sqrt{\sum_s TD_{k,s}^2}} \quad (23)$$

其中, 节点 $s \in S$, 是与节点 i 和 k 均有过直接交互的节点. 本文使用相似度 $Sim_{i,k}$ 作为节点 i 对推荐节点 k 的关系可信度 $RC_{i,k}$, 再利用式(21)得到推荐可信度 $C_{i,k}$, 同理, 若 $C_{i,k} < 0.5$, 则认为该推荐节点不可信, 直接排除. 间接推荐信任的计算如下:

$$TR_{i,j}^{in} = \sum_{k=1}^{N^{in}} \frac{C_{i,k}}{\sum_{k=1}^{N^{in}} C_{i,k}} \times TD_{k,j} \quad (24)$$

其中 N^{in} 为间接推荐节点的个数.

3.2.4 陌生推荐信任

本文将与评估节点 i 既无直接交互关系也无间接交互关系的推荐节点 k 定义为陌生推荐节点. 因无任何交互关系, 故无需考虑该类节点的关系可信度. 对于有反馈可信度的节点, 可对反馈可信度进行筛选, 若 $FC_{i,k} < 0.5$, 则认为该推荐节点不可信, 将其推荐信任丢弃; 否则认为该节点的推荐信息可信. 对于没有反馈信任度的节点, 本文利用数据统计理论中箱形图方法来剔除那些异常的推荐信任.

定义 2. 四分位数(Quartile). 统计学中分位数的一种, 将所有数值由小到大排序后分成四等份, 处于三个分割点位置的数值称为四分位数. 如 Q_1 、 Q_2 、 Q_3 被称为四分位数, Q_1 为第一四分位数, Q_2 为第二四分位数, Q_3 为第三四分位数.

箱形图识别异常值的标准为: 小于 $Q_1 - 1.5IQR$ 或大于 $Q_3 + 1.5IQR$, 其中 $IQR = Q_3 - Q_1$. 通过反馈可信度和箱形图筛选后的推荐信任集合为 $B = \{b_1, b_2, \dots, b_N^u\}$, 其中 N^u 为集合个数, 则陌生推荐信任的计算如下:

$$TR_{i,j}^u = \sum_{k=1}^{N^u} b_k / N^u \quad (25)$$

3.2.5 总体推荐信任计算

在现实生活中, 人们更相信熟人给出的推荐意见, 本文采用基于熵权的模糊层次分析法^[23]来确定三种推荐信任权重分配的最优方案.

首先, n 表示推荐信任的决策要素个数, 本文中推荐信任的决策要素分别为直接、间接、陌生推荐信任, 即 $n=3$. 用 m 表示权重分配方案个数, 用模糊数集合 $\{\tilde{1} \ \tilde{3} \ \tilde{5} \ \tilde{7} \ \tilde{9}\}$ 标识判断矩阵中的元素值, 并将层

次分析法中的两两比较改为同一决策要素下不同方案的比较, 由此得出判断矩阵 $\tilde{X}$:

$$\tilde{X} = \begin{bmatrix} \tilde{x}_{11} & \cdots & \tilde{x}_{1n} \\ \vdots & \ddots & \vdots \\ \tilde{x}_{m1} & \cdots & \tilde{x}_{mn} \end{bmatrix}.$$

用判断矩阵 $\tilde{X}$ 中的各元素乘以对应的各决策要素的模糊主观权重向量 $\tilde{W}$, 将 $\tilde{X}$ 转化为模糊判断矩阵 $\tilde{J}$:

$$\tilde{J} = \begin{bmatrix} \tilde{w}_1 \times \tilde{x}_{11} & \cdots & \tilde{w}_n \times \tilde{x}_{1n} \\ \vdots & \ddots & \vdots \\ \tilde{w}_1 \times \tilde{x}_{m1} & \cdots & \tilde{w}_n \times \tilde{x}_{mn} \end{bmatrix}.$$

其次, 通过定义置信度 α 的区间来表示三角模糊函数, 对任意 $\alpha \in [0, 1]$, 有:

$$\hat{J}_\alpha = [a_1^\alpha, a_3^\alpha] = [a_1 + (a_2 - a_1)\alpha, a_3 - (a_3 - a_2)\alpha].$$

利用三角模糊数, 我们把模糊判断矩阵 $\hat{J}$ 简化为

$$\hat{J}_\alpha = \begin{bmatrix} [a_{11L}^\alpha, a_{11R}^\alpha] & \cdots & [a_{1nL}^\alpha, a_{1nR}^\alpha] \\ \vdots & \ddots & \vdots \\ [a_{m1L}^\alpha, a_{m1R}^\alpha] & \cdots & [a_{mnL}^\alpha, a_{mnR}^\alpha] \end{bmatrix},$$

其中: $a_{ijL}^\alpha = w_{iL}^\alpha \times x_{ijL}^\alpha$, $a_{ijR}^\alpha = w_{iR}^\alpha \times x_{ijR}^\alpha$.

然后, 将模糊判断矩阵 $\hat{J}_\alpha$ 变形为非模糊判断矩阵 $\hat{A} = (\hat{a}_{ij}^\alpha)_{mn}$. 为此引入乐观指数 λ (λ 表示决策者对决策结果的乐观程度), λ 值大表示一个较高的乐观度, 其中: $\hat{a}_{ij}^\alpha = \lambda \hat{a}_{ijL}^\alpha + (1 - \lambda) \hat{a}_{ijR}^\alpha$, $\forall \lambda \in [0, 1]$. 为了克服不同维度的影响, 需将变形矩阵 $\hat{A}$ 转换成为标准矩阵 $A' = (a'_{ij})_{mn}$, 其中 $a'_{ij} = \frac{\hat{a}_{ij}^\alpha - ave_j}{dev_j}$, ave_j 是矩阵 $\hat{A}$ 中第 j 个决策要素的平均值, dev_j 是其对应的标准差.

最终根据信息熵的定义得出第 j 个决策要素的熵权:

$$e_j = -\ln(m)^{-1} \sum_{i=1}^m y_{ij} \log_2 y_{ij} \quad (26)$$

其中 $y_{ij} = a'_{ij} / \sum_{i=1}^m a'_{ij}$, 根据信息熵的计算公式得出第 j 个决策要素的权重为

$$w_j = \frac{1 - e_j}{\sum_{i=1}^m (1 - e_j)} \quad (27)$$

通过加权求和的方法来计算每一个权重分配方案的结果 $U = \sum_{i=1}^n y_{ij} w_j$, 可行方案的 U 值越大, 表示该方案的效果越好. 由此, 可以确定出最优的权重分配方案. 总推荐信任计算如下:

$$TR_{i,j} = \omega^d TR_{i,j}^d + \omega^{in} TR_{i,j}^{in} + \omega^u TR_{i,j}^u \quad (28)$$

其中, $\omega^d, \omega^{in}, \omega^u$ 分别为直接、间接、陌生推荐信任所占的权重。

3.2.6 推荐信任模型算法描述

该算法的时间复杂度为 $O(N^3)$, 与传统的推荐信任模型不同, N 不是所有推荐节点的个数, 而是直接、间接、陌生推荐节点个数中的最大值, 表明使用推荐节点分类的方法可以降低算法的时间复杂度。该算法的实现需要收集大量的推荐信任信息, 其中直接推荐信任信息很容易收集, 但间接和陌生推荐信任信息则需要在网络中至少传播两跳才能到达评估节点, 因此会带来较大的网络开销。

算法 2. 推荐信任模型算法描述。

输入: $TD_{i,k}, RT_{i,j}$

输出: $TR_{i,j}$

1. IF there is a direct interaction between k and i THEN
2. $RC_{i,k} = TD_{i,k}$;
3. computes $C_{i,k}$ using $RC_{i,k}$ and $RT_{i,j}$;
4. computes $TR_{i,j}^d$;
5. ELSE IF i and k have a common set S of interactive nodes THEN
6. computes $Sim_{i,k}$ according $TD_{i,s}$ and $TD_{k,s}$;
7. $RC_{i,k} = Sim_{i,k}$;
8. computes $C_{i,k}$ using $RC_{i,k}$ and $RT_{i,j}$;
9. computes $TR_{i,j}^{in}$;
10. ELSE
11. extract B using the box plot method and $FC_{i,k}$;
12. computes $TR_{i,j}^u$;
13. END IF
14. obtain $\omega^d, \omega^{in}, \omega^u$ by fuzzy analytic hierarchy process based on entropy weight;
15. $TR_{i,j} = \omega^d TR_{i,j}^d + \omega^{in} TR_{i,j}^{in} + \omega^u TR_{i,j}^u$;
16. RETURN $TR_{i,j}$;

3.3 总体信任模型

3.3.1 总体信任值计算

依据人类社会学, 在可信度的评估过程中, 人们主观经验的影响会更加明显, 来自他人的推荐经验仅当作参考。因此, 当自身有足够的证据判断对方是否可信时, 直接信任即为总体信任; 当无任何证据时, 推荐信任即为总体信任。 $T_{i,j}$ 计算如下:

$$T_{i,j} = \begin{cases} TD_{i,j}, & h \geq H \\ \frac{1}{1+\beta(j)} \times TD_{i,j} + \frac{\beta(j)}{1+\beta(j)} \times TR_{i,j}, & 0 < h < H \\ TR_{i,j}, & h = 0 \end{cases} \quad (29)$$

其中, h 表示两个节点直接交互的总次数, H 表示交互次数的阈值。当直接交互次数大于阈值时, 使用直接信任值作为总体信任值。 $\beta(j)$ 为 j 的实体活跃度, 它与交互节点和推荐节点的个数有关。 $\beta(j)$ 的计算公式如下:

$$\beta(j) = \frac{1}{2} [\Phi(N^k) + \Phi(N^j)] \quad (30)$$

其中 $\Phi(x) = 1 - \frac{1}{x + \delta}$, δ 为大于 0 的常数。 N^k 为推荐节点的总个数, N^j 为直接推荐节点的总个数。

3.3.2 总体信任模型算法描述

当直接交互次数介于 0 到直接交互次数阈值 H 之间时, 该算法的时间复杂度为直接信任模型与推荐信任模型算法中时间复杂度的最大值。但是空间复杂度却是二者空间复杂度之和, 因为该算法既需要存储收集到的直接信任信息, 还要存储推荐信任信息, 所需存储空间较大。而且这些信息的收集带来的网络开销也是二者网络开销之和。然而, 当直接交互次数大于交互次数阈值时, 该算法的复杂度与直接信任模型算法相同。同时, 由于不需要额外收集推荐信任信息, 极大降低了网络开销。同理, 当没有直接交互时, 该算法的复杂度等于推荐信任模型算法的复杂度。同时, 因为不需要收集直接信任信息, 网络开销也有所降低。

算法 3. 总体信任模型算法描述。

输入: $TD_{i,j}, TR_{i,j}, h, H$

输出: $T_{i,j}$

1. IF $h \geq H$ THEN
2. $T_{i,j} = TD_{i,j}$;
3. ELSE IF $h = 0$ THEN
4. $T_{i,j} = TR_{i,j}$;
5. ELSE
6. computes $\beta(j)$ by N^k and N^j ;
7. $T_{i,j} = 1/(1+\beta(j)) \times TD_{i,j} + \beta(j)/(1+\beta(j)) \times TR_{i,j}$;
8. END IF
9. RETURN $T_{i,j}$;

4 信任模型性能分析

本文使用 NetLogo 仿真平台^①分别验证了直接信任模型、推荐信任模型和总体信任模型的有效性和准确性。

① Wilensky U. NetLogo [Online]. Available: 2017. <https://ccl.northwestern.edu/netlogo>

4.1 直接信任模型预测结果对比

本实验模拟了正常节点,发起灰洞攻击节点和 on-off 攻击节点的数据转发情况,并统计了它们前 16 个周期内的行为值(300 次实验取平均值),其中前 14 个周期的行为值如表 1 所示.以此为样本,分别使用 FTI-RNBP^[21] 和 FTI-MSCGM 算法对节点未来直接信任值进行预测.

表 1 节点历史行为值记录

周期	正常节点	灰洞攻击节点	on-off 攻击节点
1	0.8486	0.2505	0.9792
2	0.9171	0.2459	0.9528
3	0.8967	0.2443	0.4378
4	0.9360	0.2397	0.4120
5	0.8434	0.2575	0.4913
6	0.8700	0.2768	0.9766
7	0.8502	0.2983	0.9452
8	0.8305	0.3023	0.9780
9	0.8887	0.2738	0.4911
10	0.9221	0.3019	0.4827
11	0.9174	0.3385	0.4773
12	0.8415	0.3266	0.9716
13	0.8818	0.3594	0.9774
14	0.9133	0.2688	0.9765

选取灰洞攻击行为值序列为例,通过对该序列进行分析和统计,获取级比序列进行波动识别,得出第 14 个周期的行为值为突变波动.因此在预测下个周期节点的信任值时应使用预测算法(2),即式(7).预测过程如下:

通过建立 SCGM(1,1)模型,可以得到拟合值 $\hat{x}^{(0)}(k)$ 和灰拟合精度值 $y(k)$,并基于 $y(k)$,进行状态划分,各状态及其范围如下: $E_1:0.5694\sim0.6631$, $E_2:0.6631\sim0.7743$, $E_3:0.7743\sim0.9264$, $E_4:0.9264\sim1.0977$.各周期的拟合精度值及对应状态见表 2.

表 2 灰拟合精度及状态划分

周期	实际值	拟合值	灰拟合精度	状态
1	0.2505	0.2282	1.0977	E_4
2	0.2459	0.2413	1.0190	E_4
3	0.2443	0.2552	0.9572	E_4
4	0.2397	0.2698	0.8884	E_3
5	0.2575	0.2854	0.9022	E_3
6	0.2768	0.3018	0.9171	E_3
7	0.2983	0.3191	0.9348	E_4
8	0.3023	0.3375	0.8946	E_3
9	0.2338	0.3569	0.6550	E_1
10	0.2754	0.3774	0.7299	E_2
11	0.3385	0.3991	0.8481	E_3
12	0.3266	0.4221	0.7737	E_2
13	0.3594	0.4463	0.8052	E_2
14	0.2688	0.4720	0.5694	E_1

根据状态划分结果,计算各步长的状态转移矩阵,其中最大步长 k 取值为 5.依据 10~14 周期,计

算各阶自相关系数并作归一化处理,得到各阶马尔可夫链权重.最终通过加权马尔可夫链预测概率 P_i 的计算来确定第 15 个周期的状态,如表 3 所示.

表 3 转移概率计算

周期	状态	步数	E_1	E_2	E_3	E_4	θ_k
14	E_1	1	0	1	0	0	0.5413
13	E_2	2	1/2	1/2	0	0	0.2595
12	E_2	3	0	1	0	0	0.1165
11	E_3	4	1/4	2/4	1/4	0	0.0590
10	E_2	5	0	0	0	0	0.0237
P_i 加权求和			0.2773	0.5752	0.1475	0	1.0000

由表 3 可知, $P_2=0.5752$ 为最大值,即下一周期的灰预测精度值最可能处于状态 E_2 ,利用式(18)在区间 $[0.6631,0.7743]$ 插值计算得:

$$\begin{aligned}\hat{y}(15) &= 0.6631 \times \frac{0.2773}{0.2773+0.1475} + \\ &\quad 0.7743 \times \frac{0.1475}{0.2773+0.1475} \\ &= 0.7017,\end{aligned}$$

由式(17)得最终预测值为

$$\begin{aligned}x^{(0)}(15) &= \hat{x}^{(0)}(15) \times \hat{y}(15) = 0.4911 \times 0.7017 \\ &= 0.3502.\end{aligned}$$

同理可得第 16 周期的预测值,如表 4 所示.预测结果表明,新算法较大幅度地提高了预测的精度,新算法的预测误差在 5%~6%左右,FTI-RNBP 算法的误差在 9%~11%左右.

表 4 预测结果对比

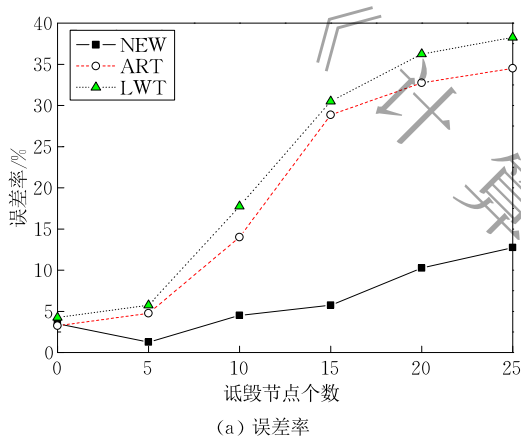
周期	真实值	拟合值	FTI-RNBP ^[21]		FTI-MSCGM	
			预测值	误差率/%	预测值	误差率/%
15	0.3314	0.4991	0.2616	10.5	0.3502	5.6
16	0.3228	0.5274	0.2989	9.9	0.3497	5.1

对于正常节点的行为值分析,可以判断没有波动值存在,预测值利用式(8)进行计算.首先获取对应的级比序列,对级比序列进行预测,预测过程同上.由于新算法的预测精度高于算法 FTI-RNBP,因此得到的级比序列值更精确,最终得出的预测结果精度也更高.

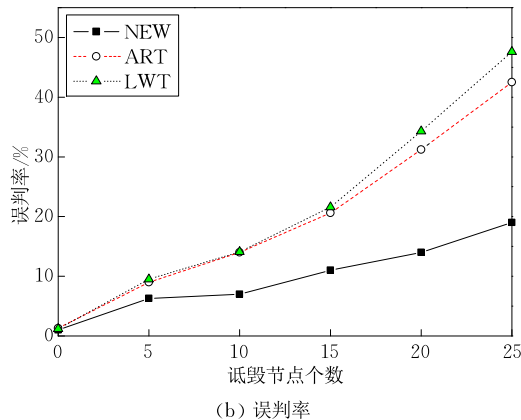
对于发起时序变换类攻击的节点,其行为周期性改变.利用算法 FTI-MSCGM,通过波动值识别和波动趋势判断,可以判断出该节点发起此类攻击.对该节点的行为值进行预测时,应将表现正常的行为监测值剔除,利用式(6)做出正确的预测.而使用算法 FTI-RNBP,则不能识别时序变换类攻击,进行预测后得到预测结果为 0.9921,认为该节点可信.但在实际中该节点此时变为恶意状态,行为值 0.4226,而造成误判.

4.2 推荐信任模型有效性验证

本实验使用两个评估指标来比较并验证推荐信任机制的有效性和准确性：(1) 误差率表示恶意节点的推荐信任值与真实信任值之间的误差；(2) 误判率表示可信节点被误判的概率，即判断出的恶意节点中可信节点所占的比例. 通过将本文提出的推荐信任模型与 ART^[17] 和 LWT^[24] 中的推荐信任模型进行对比, 设置网络节点总个数为 100, 恶意节点的丢包率为 60%, 逐渐增加诋毁节点的个数. 诋毁节点在信任推荐时, 将可信节点的信任值降低 0.5, 同时将恶意节点的信任值提高 0.5, 从而干扰推荐模型对恶意节点和可信节点信任值的判断, 影响推荐信任的准确性. 各模型的误差率和误判率随诋毁节点增多的变化趋势如图 1 所示.



(a) 误差率



(b) 误判率

图 1 诋毁节点增多对推荐信任模型的影响

如图 1(a) 所示, 新模型的误差率随着诋毁节点的增多而增长缓慢, 但模型 ART 和 LWT 的误差率骤然增加. 其中误差率大于 25% 表示模型对恶意节点的推荐信任值大于 0.5, 即认为该恶意节点误判为可信节点. 由图可知, 当诋毁节点个数超过 15 时, ART 和 LWT 中的推荐信任模型将很难对恶意节点做出正确的评估. 因为 LWT 的推荐模型简单, 对

诋毁攻击应对能力差; 而 ART 中的推荐算法主要依赖于协同过滤, 诋毁节点增多会影响其对最相似节点的选取, 导致选择的节点并不可信, 从而造成推荐信任值误差较大; 新模型在推荐可信度计算时新提出了反馈可信度的概念, 可以有效地应对诋毁节点增多对可信度计算的影响, 此外对推荐节点的分类又进一步增加了推荐信任计算的准确性.

如图 1(b) 所示, 对可信节点的误判率随诋毁节点的增多呈持续上升的趋势. 其中新模型的误判率增长较慢, 而 ART 和 LWT 的误判率在诋毁节点个数超过 10 后基本呈线性增长. 这是因为 ART 和 LWT 不能有效应对诋毁攻击, 大量诋毁节点降低了可信节点的推荐信任值, 使模型对可信节点的推荐信任值评估较低, 从而造成误判. 这也再次证明了新模型中推荐信任模型的有效性和准确性.

4.3 总体信任模型性能对比

本实验模拟实现了具有灰洞攻击、on-off 攻击和诋毁攻击的网络, 将新模型与信任模型 ART^[17] 和 CAT^[18] 进行仿真对比, 证明了新模型的有效性. 使用两个指标来衡量模型应对恶意攻击的能力: (1) 恶意节点检出率, 即识别出的真实恶意节点数量与总恶意节点数量的比值; (2) 恶意节点识别精确度, 即识别出的真实恶意节点数量与识别出的恶意节点数量的比值. 信任模型参数、仿真参数和情景设置如表 5、表 6、表 7 所示.

表 5 信任模型省缺参数

参数	值	参数	值
交互次数调节因子 a	4	惩罚因子 p	0.25
偏离阈值 d	0.2	直接交互次数阈值 H	1000
奖励因子 r	0.15	活跃度调节因子 δ	0.8

表 6 仿真参数

参数	值
仿真时间	2000 s
节点个数	100
恶意节点	10, 20, 30

表 7 情景设置

攻击类型	描述
灰洞攻击	恶意丢包 50%
on-off 攻击	一段时间内正常转发数据包, 然后再恶意丢包 50%, 交替进行, 周期为 40 s
诋毁攻击	信任推荐时, 降低诚信节点的信任值(减 0.5), 同时提高不良节点的信任值(加 0.5)

各模型面对不同攻击时, 对恶意节点检出率的对比如图 2 所示. 在灰洞攻击下, 三种模型的检出率

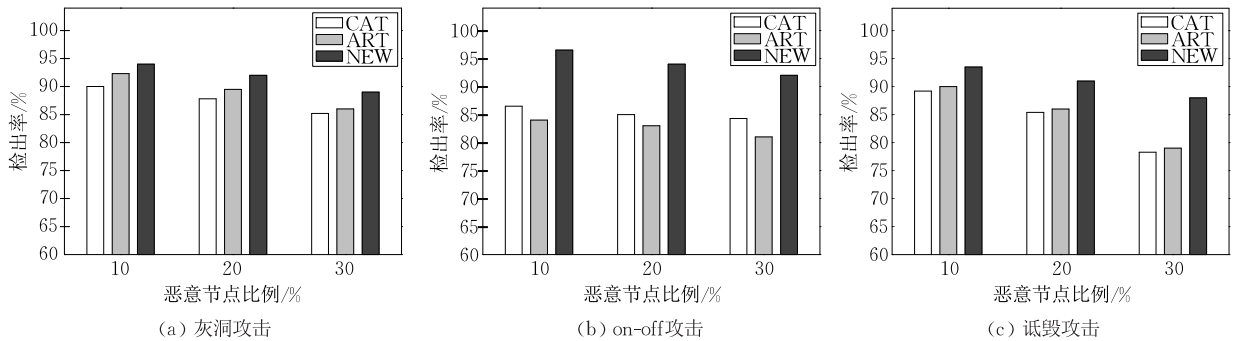


图 2 不同攻击下各模型恶意节点检出率对比

差别不大,如图 2(a)所示. 这表明对于简单的恶意丢包类攻击,三种模型都可以有效地应对. 信任计算越精确,恶意节点的检出率就越高. CAT 模型可以通过基于 MDS(恶意行为检测机制)的评估模型检测出节点的恶意丢包行为;ART 模型可以通过 DS 证据理论将多个证据组合起来,对节点的数据信任和功能信任进行适当的评估. 但 CAT 模型侧重于对车辆行为的检测,忽略了对其它车辆推荐信任信息的收集计算;ART 模型则依赖于间接信任,对主观的直接信任考虑不足;而新模型综合了直接信任和间接信任,使得信任计算的精确度更高,因此新模型的检出率略高于其它两种模型.

如图 2(b)所示,在应对 on-off 攻击时,CAT 模型和 ART 模型的检出率较低,因为它们并未提出有效的机制去应对时序变换类攻击,很难识别出发起该类攻击的恶意节点. 此外,由于 CAT 模型侧重于检测节点恶意行为,对历史交互记录依赖较小,当具有 on-off 攻击的节点发生恶意丢包时,可能会错误地认为该攻击为灰洞攻击,从而判断该节点为恶意节点,而 ART 模型侧重于通过收集节点的历史证据,很容易对该类节点做出误判,所以 CAT 模型的检出率略高于 ART 模型. 新模型在直接信任计算过程中,基于识别出的波动值,提出了波动趋势

的概念,专门针对时序变换类攻击,能有效地识别出节点的 on-off 攻击,因此检出率较高,且保持在 90%以上.

如图 2(c)所示,当诋毁节点较少时,模型 CAT 和 ART 与新模型的检出率相差不大. 因为二者都有应对诋毁攻击的机制,CAT 提出了 k 均值聚类的方法,ART 提出了基于协作过滤的推荐策略. 但随着诋毁节点的增多,这两种方法应对诋毁攻击的能力逐渐下降,检出率快速降低. 而新模型在推荐信任模型中提出了反馈可信度的概念,利用交互结果检验节点推荐信任的准确性,不受恶意节点增多的影响,能有效地避免诋毁攻击的影响,对节点属性做出正确的评估. 因此,新模型在诋毁攻击下的检出率与灰洞攻击下的检出率基本相同,受诋毁攻击影响较小.

各模型面对不同攻击时对恶意节点识别精确度的变化如图 3 所示. 在面对灰洞攻击时,三种模型对恶意节点的识别精确度基本相似,如图 3(a)所示. 因为此时网络中没有传播任何虚假的信任推荐意见,各模型在恶意节点识别过程中没有受到虚假信息的影响,对恶意节点的识别精确度取决于信任预测算法的精确度,新模型综合直接与推荐信任模型,预测精确度优于 ART 模型和 CAT 模型,因此新模型的识别精确度略高.

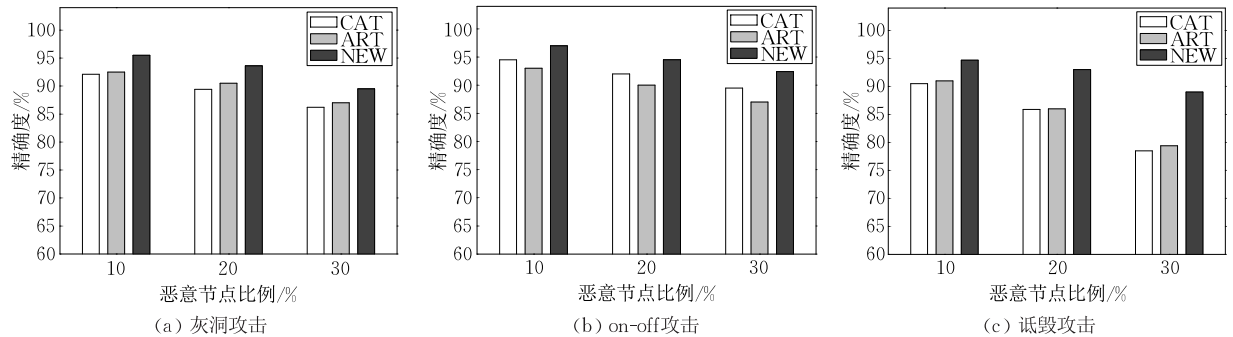


图 3 不同攻击下各模型恶意节点识别精确度对比

如图 3(b)所示,虽然 ART 模型和 CAT 模型对 on-off 攻击的节点检出率不高,但因为没有外界虚假信息的影响,对恶意节点的识别精确度依然保持较高的水平.此时精确度的高低主要取决于对恶意节点的检出率,检测出的真实恶意节点越多,精确度越高,因此,新模型的对恶意节点识别的精确度最高,CAT 模型次之,ART 模型最低,但差距并不明显.

在面对诋毁攻击时,随着恶意节点的增多,会发生更为严重的共谋攻击,模型 CAT 和 ART 由于推荐信任机制并不完善,对节点可信度的误判机率增大,对恶意节点的识别精确度随之迅速下降.而本文提出的新模型能有效地应对诋毁攻击,过滤不良推荐,降低对节点可信度的误判,因此精确度仍然保持较高的水平,如图 3(c)所示.

5 基于信任的安全组播路由协议

以信任计算模型为基础,借鉴路由策略中网络结构的优点,本节提出一种适用于 VANETs 的基于信任的安全组播路由协议 TSMRP.该协议不仅能适应 VANETs 网络中车辆的高速移动,而且能通过信任计算排除恶意车辆节点,保证网内数据信息安全可靠地传输.

5.1 TSMRP 路由发现

5.1.1 路由请求

源节点首先广播路由请求报文 JOIN RREQ 并试图与目标节点建立安全可靠的信息传播路径.报文 JOIN RREQ 被源节点周期性向全网广播,用于刷新组成员的路由信息,避免因路由失效造成数据传输的中断.当中间节点收到 JOIN RREQ 后,利用消息缓存表和恶意节点表,判断该报文是否是重复报文以及上游节点是否是恶意节点;若该报文不是重复报文且上游节点不是恶意节点,将 JOIN RREQ 加入消息缓存表,并将 JOIN RREQ 中的路由信息存入节点的路由表中,据此建立反向路由,否则将路由请求报文丢弃;最后将 JOIN RREQ 中的上一跳字段更新后执行广播操作.

5.1.2 路由回复

接收节点收到请求报文后,立即创建路由回复报文 JOIN TABLE,然后将其广播给所有邻居节点.节点收到 JOIN TABLE 后,首先检查自己的恶意节点表.若判断出上游节点为恶意节点,则将该报文丢弃;若为正常节点,则将自身地址与路由回复报文中的下一跳字段地址匹配,若匹配失败,将 JOIN

TABLE 丢弃;若匹配成功,则证明自己是转发组节点,需要更新 JOIN TABLE 的下一跳信息,并将其继续广播给邻居节点.

JOIN TABLE 由转发组节点广播,直到通过最短的可信路径到达组播源.此过程建立(或更新)从源到接收节点的路由,并构建转发组.换句话说,在建立的从源节点到接收节点的网格结构中,每个节点都是可信的.

5.2 TSMRP 路由维护

因存在路由路径失效问题(如节点移动性导致路由路径断裂),为了保证路由的有效性和新鲜性,源节点需要周期性地广播路由维护报文,以修复过期或断裂的路由路径.

5.2.1 路由效率提高机制

为了进一步提高路由效率,在路由维护阶段本文提出两种改进机制.

(1)转发组节点复用机制.该机制适用于网格建立后,有新节点加入多播组的情况.每个转发组节点可能会收到来自多个源节点广播的 JOIN RREQ 报文,转发组节点都会执行转发操作,可考虑减少转发组节点的总体数量以减少路由维护报文的洪泛.

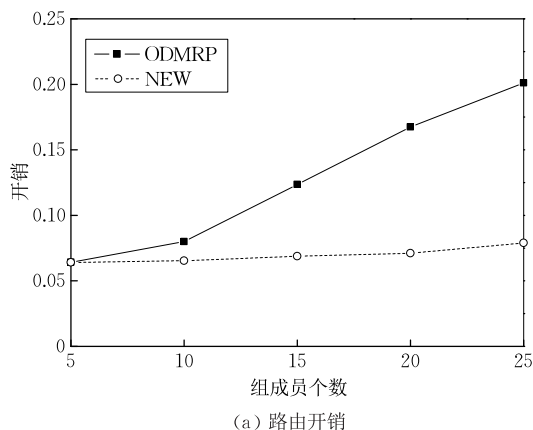
首先需要记录转发组节点的数目(可利用 JOIN RREQ 的 FC 字段),中间转发节点若为转发组节点,则将该字段的值加 1.随后广播更新的 JOIN RREQ 报文直到新节点收到该报文.新节点并不立即对最先接收到的 JOIN RREQ 报文进行回复,而是等待一段时间,对比多个 JOIN RREQ 报文中 FC 字段的值,对最大 FC 值的维护报文执行回复.选取的路径并不一定是最短路径,但转发组节点数量的减少却可以阻止大量 JOIN RREQ 报文的洪泛,从而大幅度降低网络开销,提高路由效率.

(2)边缘路径删除机制.该机制主要用于维护已经加入多播组的接收节点到源节点的路由.网格结构之所以具有较强的鲁棒性,是因为当链路断开时,数据包可以通过冗余链路继续进行数据传输,而冗余链路又依赖于相邻的转发组节点.对于网格中的其中一条路径,所有的转发组节点除了自己的上游和下游节点外,周围并没有或者只有极少量的组成员节点,本文定义该路径为边缘路径.如果这些路径发生了链路中断,则很难通过冗余链路进行直接修复,从而影响数据包的传递率.因此,删除边缘路径,选取具有较多组成员邻居的节点所在的路由作为最终路由的方法,不仅可以增强网格结构的稳定性,还可以进一步减少转发组节点的个数,提高路由

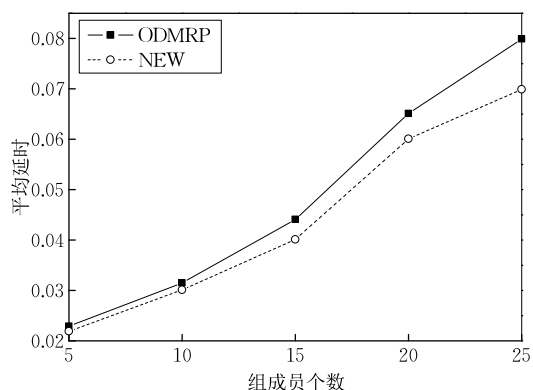
效率. 为了实现上述过程, 本文在 JOIN RREQ 和路由表中加入新的表项 GNC (Group Neighbor Counts) 和 GN (Group Neighbor), 分别代表该路径上组成员邻居的总个数和某个节点周围的组成员邻居个数. 源节点周期性的洪泛 JOIN RREQ 进行路由维护时, 初始化 GNC 的值为 0. 中间节点在收到该 JOIN RREQ 后, 若是转发组节点, 从 Routing Table 中取出 GN 的值, 与 GNC 相加, 作为新的 GNC 存入 JOIN RREQ, 然后继续转发. 当接收节点收到最先到达的 JOIN RREQ 后, 判断其中的 GFC 的值, 若 GFC 的值大于阈值 β (β 为大于 0 的常数, 下文实验取 $\beta=2$), 则认为该路径不是边缘路径, 然后直接进行路由回复; 若 GFC 的值小于阈值 β , 则认为该路径是边缘路径, 延时一段时间 τ 后, 选取具有最大 GFC 的 JOIN RREQ 进行路由回复.

5.2.2 优化机制性能分析

本节使用 NS-2 仿真平台^①对协议性能进行仿真, 结果如图 4 所示. 图 4(a) 表示随着组成员个数的增多, ODMRP 协议由于大量请求报文的洪泛造成路由开销快速增大, 而新协议由于控制了转发组节点个数, 开销增长缓慢. 图 4(b) 所示的是平均延时随组成员节点增多的变化趋势. 虽然新协议建立



(a) 路由开销



(b) 平均延时

图 4 改进路由算法与 ODMRP 性能对比

的不是最短路径, 增加了路由的跳数, 但减少了路由过程中的链路阻塞, 降低了因链路阻塞导致的缓冲队列延时, 使得路由过程中总的平均延时反而低于 ODMRP 协议. 因此转发组节点复用机制和边缘路径删除机制的提出使新协议进一步提高了路由的整体效率.

6 实验结果及分析

本节的实验仿真使用 SUMO 仿真平台^[25]和 NS-2 仿真平台共同完成. SUMO 能够实现车辆在特定地图中的仿真运行, 负责模拟仿真场景的构建, 包括道路模型的设计和车辆移动模型的建立, 为下一步协议的网络仿真提供了相对真实的仿真场景. 随后将 SUMO 中车辆仿真模型的输出文件 (activity、mobility 和 config 文件) 嵌入 NS-2 中进而对协议的性能进行仿真和验证.

本文使用以下四个评估指标对新协议 TSMRP 与 ODMRP 以及两种经信任扩展的协议 ART-ODMRP^[17]和 CAT-ODMRP^[18]进行仿真对比: (1) 包投递率 (Packet Delivery Ratio, PDR): 目标节点接收到的数据包与源节点发送的数据包的比值关系; (2) 路由开销 (Routing Overhead): 网络通信过程中控制包数量与数据包数量的比值; (3) 端到端平均延时 (Average End-to-End Latency): 数据流中的数据包通过网络从源节点传输到目标节点所花费的平均时间; (4) 单位传包量 (Byte Sent per Byte Delivered, BSBD): 网络中转发节点传递的数据包总数量与接收的数据包总数量的比值.

6.1 场景构建和参数设置

VANETs 仿真中模拟场景的创建需真实地反映道路交通环境, 本文使用 OpenStreetMap 地图编辑工具, 手动选择青岛市市北区的部分地图作为仿真的道路模型, 如图 5 所示.

使用 JOSM 软件对地图进行调整, 删去地图中的各种建筑物, 只保留道路信息, 并对地图中的道路进行微调, 以适应于协议的仿真. 随后通过 NETCONVERT 工具将地图数据 OSM 文件转化为 SUMO 的道路模型文件. JOSM 可以获取每个路口所在位置的 id, 实验中在十字路口和三叉路口处加入交通灯. 最终可使用 SUMO-GUI 查看生成的道路地图, 如图 6 所示.

① A discrete event simulator NS-2 [Online]. Available: 2017. <https://www.isi.edu/nsnam/ns/>



图 5 OpenStreetMap 中青岛市市北区部分地图

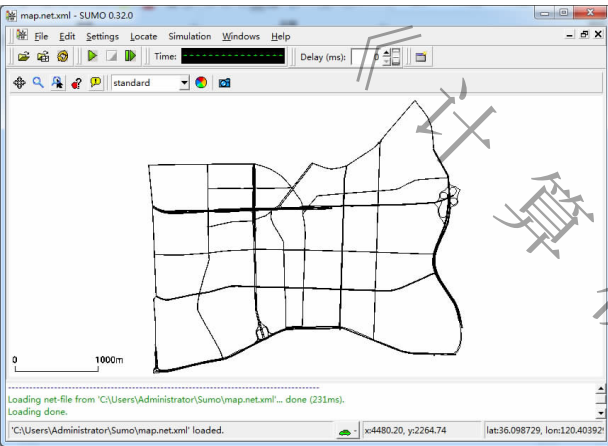


图 6 SUMO 中道路模型

道路模型创建成功后,还需要使用 SUMO 产生车辆移动模型.为了保证模拟的真实性,本文随机设置车辆在地图中的起点和终点.车辆的仿真参数与地图参数设置如表 8 所示.

表 8 SUMO 仿真参数

参数	值	参数	值
地图区域	青岛市市北区	车辆个数	100
仿真区域大小	2500 m×3500 m	车辆长度	5 m
路口个数	47	车辆最大速度	5~30 m/s
红绿灯个数	36	车辆加速最大加速度	0.8 m/s ²
道路条数	15	车辆减速最大加速度	4.5 m/s ²

仿真场景为 SUMO 中的车辆仿真模型,设置每次仿真的时间为 3000 s,取 30 次实验结果的平均值作为模拟的最终结果.此外,经信任扩展后的三个协议均可有效地应对灰洞攻击,因此本节模拟车辆发起诋毁攻击和 on-off 攻击的情景.各组测试参数如表 9 所示.

表 9 各组测试参数列表

测试组号	源车辆个数	目的车辆个数	最大速度/(m·s ⁻¹)	恶意车辆个数
1	10	10	10	0,10,20,30,40,50
2	10	10	5,10,15,20,25,30	20

6.2 恶意车辆数量的改变对协议性能的影响

设定网络中车辆最大移动速度为 10 m/s,观察恶意车辆数量的改变对协议性能的影响.

如图 7(a)所示,ODMRP 协议的包投递率随恶意车辆的增多而迅速下降;经信任扩展的 TSMRP、ART-ODMRP 和 CAT-ODMRP 协议,能识别恶意车辆,在恶意车辆个数不超过 30 时,包投递率缓慢降低,且不低于 90%.当恶意车辆超过 30 时,恶意车辆数目对网络的影响变大,各协议的识别效果减弱,包投递率下降速率加大. ART-ODMRP 和 CAT-ODMRP 协议随恶意车辆的增多,应对 on-off 攻击和诋毁攻击的能力迅速下降,使得它们与 TSMRP 协议的投递率差距逐渐增大.而 TSMRP 利用基于波动的灰色马尔可夫预测模型,考虑了车辆节点历史行为和当前行为值的变化,可以有效地应对 on-off 攻击;在计算推荐信任时利用推荐可信度、相似度和箱形图,有效过滤不良推荐来应对诋毁攻击,故投递率最高.

路由开销的变化如图 7(b)所示,随着恶意车辆的增多,增加了网络的路由发现和维护开销,因此各协议的开销逐渐上升且增长幅度逐渐增大.当没有恶意车辆时,基于信任的路由协议 ART-ODMRP 和 CAT-ODMRP 仍然执行信任建模策略,相较 ODMRP 开销略微增加.然而,随着恶意车辆的进一步增多,ODMRP 接收到的数据包减少,用于路由维护的控制包增多,开销骤增;而加入信任模型后,降低了开销,因此当恶意车辆个数超过 10 时,二者的开销低于 ODMRP.新协议 TSMRP 在没有恶意车辆时也需要收集信任信息,增加了额外的开销,但是由于 TSMRP 在路由维护过程中利用转发组节点复用机制和边缘路径删除机制,降低了路由维护过程中的开销,因此在没有恶意车辆时开销略低于 ODMRP,而且随恶意车辆的增多开销一直低于其它三种协议.

如图 7(c)所示,TSMRP、ART-ODMRP 和 CAT-ODMRP 协议在可信路由路径的建立过程中只选择可信车辆节点,规避恶意车辆节点,导致路径变长,端到端平均延时大于 ODMRP 协议.当恶意车辆个

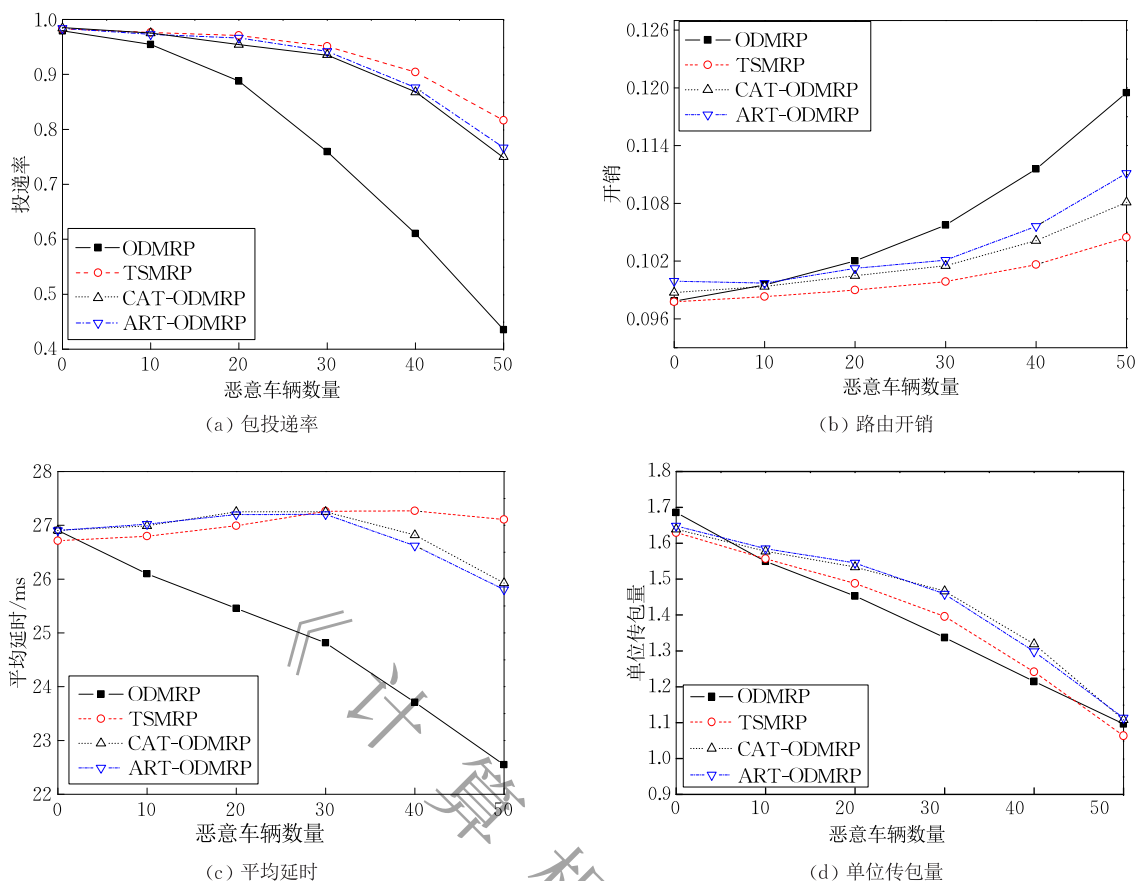


图 7 恶意车辆数量的改变对协议性能的影响

数小于 30 时,随恶意车辆的增多,为了排除恶意车辆,建立安全路由,路由跳数增加,平均延时缓慢增大.由于 TSMRP 在路由维护过程中降低了开销,减少了路由过程中的链路阻塞,在一定程度上避免了链路阻塞导致的缓冲队列延时,因此平均延时较低;当恶意车辆数大于 30 时,ART-ODMRP 和 CAT-ODMRP 协议识别恶意车辆的能力大幅度下降.由于网络中存在大量的恶意车辆,数据包在传递到较远目的车辆的过程中更容易被恶意车辆丢弃,从而无法到达.而数据包传递到距离较近的车辆的的过程中,经过的中间车辆较少,受恶意车辆影响不大,因此大部分数据包只能到达距自己较近的车辆的,导致平均延时降低.而当恶意车辆数大于 30 时, TSMRP 协议对恶意车辆的识别能力并没有明显的下降,因此平均延时仍然保持缓慢上升的趋势.

单位传包量反映了路由协议对网络流量的控制能力, TSMRP、ART-ODMRP 和 CAT-ODMRP 协议需要收集信任信息,会产生额外的流量,因此单位传包量大于 ODMRP. 随恶意车辆的增加,大量的数据包被丢弃,减少了数据包在网络中的传播,在一定程度上降低了网络流量,导致单位传包量降低,如图 7(d)所示. TSMRP 在信任计算时若有足够的直接

交互则不需要收集推荐信任信息,在一定程度上减少了网络流量,同时减少转发组节点的个数也能进一步控制了流量,因此与其它协议相比, TSMRP 能更好地控制网络流量.

6.3 车辆最大移动速度的改变对协议性能的影响

使用真实道路模型为仿真场景,保持恶意车辆个数不变,改变车辆的最大移动速度,来探究协议能否有效地适应车辆的快速移动性. 我们认为当车速超过 60 km/h (即 16.7 m/s) 时为快速移动.

如图 8(a)所示,各协议的包投递率并不随速度的增加严格地递增或递减,但整体上呈下降趋势. 当速度超过 60 km/h 时,由于各协议都属于网格结构,且具有较高的鲁棒性,投递率并没有骤然下降,仍保持在 95% 以上. 但是车辆的高速移动增加了车辆收集信任信息的难度,可能导致信任信息收集不足,因此投递率逐渐降低且降低幅度逐渐增大. 其中, CAT-ODMRP 协议中的信任模型侧重于通过收集周边信息来检测恶意车辆,当周边车辆移动到检测车辆范围外后,可能导致信任信息收集不足,而该协议没有推荐信任收集模块,因此受车辆快速移动影响最大,投递率低于 TSMRP 和 ART-ODMRP 协议; ART-ODMRP 协议则主要依赖于间接信任,对

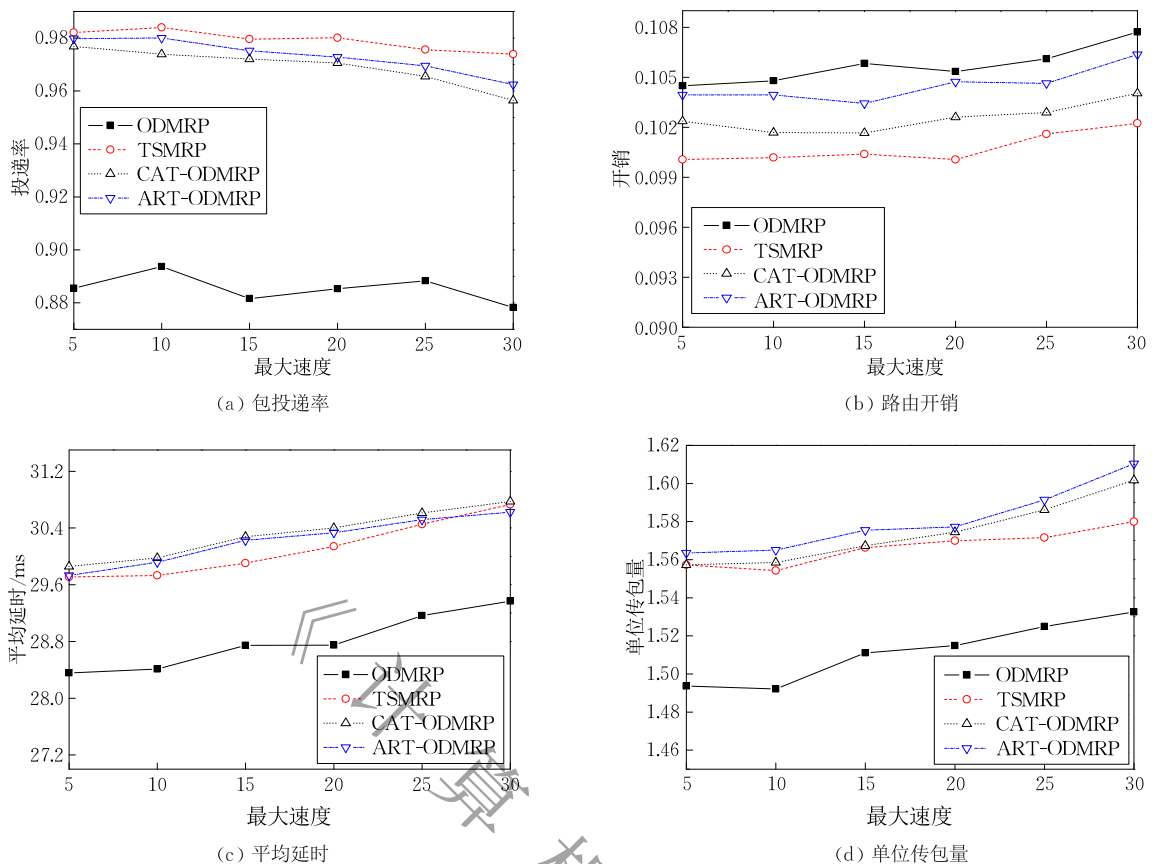


图 8 车辆最大移动速度的改变对协议性能的影响

直接信任的计算考虑不足,当间接信任模型受车辆速度影响不能准确预测信任值时,无法通过直接信任来进行弥补,因此投递率不如 TSMRP. TSMRP 全面地考虑了直接与间接信任模型,且预测精确度较高,对信任信息的收集较为全面,受车辆快速移动的影响较小,因此投递率最高.

路由开销随速度的变化情况如图 8(b)所示.随着速度的增大网络拓扑结构将频繁变化,网络更易发生断路,导致用于路由维护的开销增大.但网络结构具有较高的稳定性,因此车辆的快速移动并没有对开销产生明显的影响,开销增长幅度并不大.其中 ODMRP 协议受恶意车辆的影响开销最大;ART-ODMRP 由于采用了 D-S 证据理论,需要收集足够的信任信息,而且信任模型的设计较 CAT-ODMRP 协议复杂,因此开销比 CAT-ODMRP 大;TSMRP 协议在路由维护过程中设计了降低开销的机制,开销最低.

随着速度的增加,网络拓扑将频繁变化,导致路由由链路断开,在重新建立路径之前,数据包将一直在缓冲队列中等待,因此端到端平均延时逐渐增大.当车速超过 60 km/h 后,平均延时的增长速度加快,但是由于网络结构较高的鲁棒性,平均延时

的增长并不因车辆的快速移动而骤增,反而增长相对平稳.如图 8(c)所示,随着速度的增加,ART-ODMRP 和 CAT-ODMRP 对恶意车辆的识别能力降低,导致网络中不能识别的恶意车辆增多,这在一定程度上降低了平均延时,因此整体平均延时趋于平缓;而 TSMRP 对恶意车辆的识别能力并没有明显下降,因此平均延时持续上升并逐渐超过 ART-ODMRP 和 CAT-ODMRP 协议.

如图 8(d)所示,单位传包量都随速度的增加而逐渐递增,且增长相对平缓,并没有因车辆的快速移动而显著增加. TSMRP 对网络流量的控制由于 ART-ODMRP 和 CAT-ODMRP 协议,因此单位传包量比二者都低.

7 结束语

VANETs 是智能交通系统的重要组成部分,在提高道路安全和交通运输生产率方面发挥着重要作用.但由于缺乏公钥基础设施, VANETs 更容易遭受恶意车辆节点的恶意攻击.本文以信任计算模型为基础,提出了一种抗多种攻击的轻量级安全组播路由协议 TSMRP.该协议经信任扩展,在路由选择

过程中,中间转发节点利用自己的直接历史交互经验或其它节点的推荐信息,对它的潜在下一跳节点进行信任值的评估,从而选取可信的下一跳作为其中继.实验结果表明,新协议可在有恶意节点存在的VANETs网络中建立安全可靠的信息传输路径,以此来抵抗各种恶意攻击,保持较高的数据包投递率、较低的路由开销和端到端平均延时.

虽然基于信任管理的安全路由协议具有较大的优越性,但仍有一些潜在的问题亟需解决,如多条可信路径冲突问题,自私车辆节点激励问题等.因此,在以后的工作中将继续对该类协议作深入的研究,尽可能地完善安全路由机制.

致 谢 在此,向本文编辑和对本文提出宝贵意见的各位审稿专家表示衷心的感谢!

参 考 文 献

- [1] Cooper C, Franklin D, Ros M. A comparative survey of VANET clustering techniques. *IEEE Communications Surveys and Tutorials*, 2017, 19(1): 657-681
- [2] Bitam S, Mellouk A, Zeadally S. Bio-inspired routing algorithms survey for vehicular ad hoc networks. *IEEE Communications Surveys and Tutorials*, 2015, 17(2): 843-867
- [3] Cheng J J, Cheng J L, Zhou M C, et al. Routing in Internet of vehicles: A review. *IEEE Transactions on Intelligent Transportation Systems*, 2015, 16(5): 2339-2352
- [4] Vijayakumar P, Azees M, Kannan A, et al. Dual authentication and key management techniques for secure data transmission in vehicular ad hoc networks. *IEEE Transactions on Intelligent Transportation Systems*, 2016, 17(4): 1015-1028
- [5] Sugumar R, Rengarajan A, Jayakumar C. Trust based authentication technique for cluster based Vehicular Ad Hoc Networks (VANET). *Wireless Networks*, 2018, 24(2): 373-382
- [6] Yao J P, Feng S L, Zhou X Y. Secure routing in multihop wireless ad-hoc networks with decode-and-forward relaying. *IEEE Transactions on Communications*, 2016, 64(2): 753-764
- [7] Sharef B T, Alsaqour R A, Ismail M. Vehicular communication ad hoc routing protocols—A survey. *Journal of Network & Computer Applications*, 2014, 40(1): 363-396
- [8] Togou M A, Hafid A, Khoukhi L. SCRIP: Stable CDS-based routing protocol for urban vehicular ad hoc networks. *IEEE Transactions on Intelligent Transportation Systems*, 2016, 17(5): 1298-1307
- [9] Wang M, Shan H, Luan T H, et al. Asymptotic throughput capacity analysis of VANETs exploiting mobility diversity. *IEEE Transactions on Vehicular Technology*, 2015, 64(9): 4187-4202

- [10] Lin D, Kang J, Anna S, et al. MoZo: A moving zone based routing protocol using pure V2V communication in VANETs. *IEEE Transactions on Mobile Computing*, 2017, 16(5): 1357-1370
- [11] Wu C, Ji Y, Liu F, et al. Toward practical and intelligent routing in vehicular ad hoc networks. *IEEE Transactions on Vehicular Technology*, 2015, 64(12): 5503-5519
- [12] Zhang X M, Chen K H, Cao X L, et al. A street-centric routing protocol based on microtopology in vehicular ad hoc networks. *IEEE Transactions on Vehicular Technology*, 2016, 65(7): 5680-5694
- [13] Zhu L, Li C, Li B B, et al. Geographic routing in multilevel scenarios of vehicular ad hoc networks. *IEEE Transactions on Vehicular Technology*, 2016, 65(9): 7740-7753
- [14] Yao L, Wang J, Wang X, et al. V2X routing in a VANET based on the hidden Markov model. *IEEE Transactions on Intelligent Transportation Systems*, 2018, 19(3): 889-899
- [15] Kerrache C A, Calafate C T, Cano J C, et al. Trust management for vehicular networks: An adversary-oriented overview. *IEEE Access*, 2016, 4: 9293-9307
- [16] Eiza M H, Owens T, Ni Q. Secure and robust multi-constrained QoS aware routing algorithm for VANETs. *IEEE Transactions on Dependable and Secure Computing*, 2016, 13(1): 32-45
- [17] Li W J, Song H B. ART: An attack-resistant trust management scheme for securing vehicular ad hoc networks. *IEEE Transactions on Intelligent Transportation Systems*, 2016, 17(4): 960-969
- [18] Dietzel S, Petit J, Heijenk G, et al. Graph-based metrics for insider attack detection in VANET multi-hop data dissemination protocols. *IEEE Transactions on Vehicular Technology*, 2013, 62(4): 1505-1518
- [19] Krishna T, Raghu V, Barnwal R P, et al. CAT: Consensus-assisted trust estimation of MDS-equipped collaborators in vehicular ad-hoc network. *Vehicular Communications*, 2015, 2(3): 150-157
- [20] Marmol F G, Martinez P G. TRIP: A trust and reputation infrastructure-based proposal for vehicular ad hoc networks. *Journal of Network and Computer Applications*, 2012, 35(3): 934-941
- [21] Xia Nu, Li Wei, Luo Jun-Zhou, et al. A Routing node behavior prediction algorithm based on fluctuation type identification. *Chinese Journal of Computers*, 2014, 37(2): 326-334(in Chinese)
(夏怒,李伟,罗军舟等.一种基于波动类型识别的路由节点行为预测算法. *计算机学报*, 2014, 37(2): 326-334)
- [22] Xia H, Li Z, Zheng Y, et al. A novel light-weight subjective trust inference framework in MANETs. *IEEE Transactions on Sustainable Computing*, 2018, DOI: 10.1109/TSUSC.2018.2817219
- [23] Xia H, Yu J, Pan Z, et al. Applying trust enhancements to reactive routing protocols in mobile ad hoc networks. *Wireless Networks*, 2016, 22(7): 2239-2257

- [24] Wang B, Chen X, Chang W. A light-weight trust-based QoS routing algorithm for ad hoc networks. *Pervasive and Mobile Computing*, 2014, 13(4): 164-180
- [25] Behrisch M, Bieker L, Erdmann J, Krajzewicz D. SUMO-

simulation of urban mobility—An overview//*Proceedings of the 3rd International Conference on Advances in System Simulation*. Barcelona, Spain, 2011: 63-68



XIA Hui, Ph. D. , associate professor. His research interests include Internet of Things, wireless network, information security and trust computing.

ZHANG San-Shun, M. S. candidate. His research interests include Internet of Things, wireless network and information security.

SUN Yun-Chuan, Ph. D. , professor. His research

interests include data science, Internet of Things, and information security.

XIAO Fu, Ph. D. , professor. His research interests include Internet of Things, network communication protocols and algorithms.

LI Ye, Ph. D. , professor. His research interests include multimedia signal processing and wireless communication.

CHENG Xiu-Zhen, Ph. D. , professor. Her research interests include wireless network and communication, Internet of Things, privacy protection and distributed algorithms.

Background

As people are increasingly demanding to improve the efficiency and safety of transportation systems, which stimulated the integration of wireless network technology and vehicles, resulting in the formation of vehicular ad hoc networks (i. e. , VANETs). However, due to the increase of relying on communication, information and control technologies, VANETs are more vulnerable to security attacks from malicious vehicles. The focus of this paper is routing security problem, which aims at establishing a secure and reliable routing protocol to identify and suppress the vehicles with malicious attacks in VANETs.

At present, researchers have proposed some solutions to the problem of routing security. These solutions can be divided into two main categories, cryptography-based countermeasures and trust-based countermeasures. The latter is preferable to the former when dealing with internal malicious node attacks. However, those researchers have not comprehensively investigated how to manage trust in VANETs in a holistic manner. And few of them incorporate the concept of trust into the routing protocol designing process.

To tackle the routing security issue, this paper proposes an attack-resistant lightweight Trust-based Secure Multicast Routing Protocol (i. e. , TSMRP) based on a trust computing model. First, the paper constructs a highly efficient trust model, which integrates two decision factors, as the direct trust and the recommendation trust, to derive the overall trust value of a vehicular node. The grey Markov prediction algorithm based on fluctuation type recognition is utilized to

accurately calculate the direct trust, whereas the fuzzy analytic hierarchy process based on entropy is used to obtain the recommendation trust. Second, a trust-based secure multicast protocol is designed to effectively identify and suppress malicious vehicles during route establishment phase and maintenance phase, thereby developing communication paths of high security and reliability. Finally, compared with the other relevant protocols, the detailed experimental results show that TSMRP not only can deal with various types of malicious attacks (e. g. , grey-hole attacks, bad-mouthing attacks, and on-off attacks), but can also effectively improve the accuracy of malicious vehicle identification and the packet delivery ratio. In addition, the routing overhead, the average end-to-end latency and the byte sent per byte delivered are reduced.

This research is supported by the National Natural Science Foundation of China under Grant Nos. 61872205, 61371185, the Key Program of the National Natural Science Foundation of China under Grant No. 61832012, the National Key Research and Development Program of China under Grant No. 2018YFB0803400, the Recruitment Program of Global Experts, the Project of Shandong Province Higher Educational Science and Technology Program under Grant No. J16LN06, the Source Innovation Program of Qingdao under Grant No. 18-2-2-56-jch, the State Foundation of China for Studying Abroad to Visit the United States as a 'Visiting Scholar' and the Open Research Fund from Shandong Provincial Key Laboratory of Computer Networks under Grant No. SDKLCN-2018-07.