

· CTCIS 2016 推荐论文 ·

DOI:10.15961/j.jsuese.2017.01.027

基于自扩展时间窗的告警多级聚合与关联方法

李洪成, 吴晓平

(海军工程大学 信息安全系, 湖北 武汉 430033)

摘要:针对传统告警聚合与关联方法在合理性和准确性上的不足,提出了基于多级划分思想的告警聚合方法和基于马尔可夫链模型的告警关联方法。首先,使用入侵检测消息交换格式来描述网络告警,利用告警的时序接近关系进行时间窗口的自动扩展,将时间间隔小于预设阈值的告警划分到同一个时间窗内;进而,分别根据攻击类型、时间窗口、子网掩码、IP地址和端口信息依次划分告警,利用属性匹配方法进行子网级、主机级和服务级聚合,有效聚合攻击者利用同一路由器、傀儡主机或服务端口实施攻击而产生的相似告警;在此基础上,利用1阶马尔可夫链模型生成告警关联图,将攻击类型间的条件转移概率作为关联图的有向边,并利用告警的时序紧邻关系计算出攻击类型间的转移概率。实验中,利用入侵检测系统 Snort 的最严格模式处理 DARPA2000 流量数据,得到 LL-DoS1.0 攻击场景所对应的入侵告警集合;利用本文方法对集合中的5类告警进行聚合和关联,通过参数寻优得到自扩展时间窗口最理想的间隔阈值,使得告警多级聚合结果能够有效精简告警,并与告警源IP和源端口的分布情况一致;通过比较告警关联结果与攻击场景的官方描述来计算告警关联的准确率。与传统方法进行对比,本文方法的告警关联准确率为97.94%,比传统方法提高了2.29%。

关键词:攻击检测;告警聚合;自扩展时间窗口;多级划分;马尔可夫模型

中图分类号:TP309.7

文献标志码:A

文章编号:2096-3246(2017)01-0206-07

Multistage Aggregation and Correlation for Network Alerts Based on Self-extending Time Windows

LI Hongcheng, WU Xiaoping

(Dept. of Info. Security, Naval Univ. of Eng., Wuhan 430033, China)

Abstract: In order to deal with the shortages of traditional alerts aggregation and correlation methods on rationality and accuracy, an aggregation method based on multistage division and a correlation method based on Markov chains model were presented. Firstly, the network alerts were described by intrusion detection message exchange format. If the time internals of alerts were shorter than the predefined threshold, the alerts would be divided into the same time window, and the time windows were extended automatically based on the temporal relationship of alerts. Then, the alerts were divided respectively according to the attributes of attack types, time windows, subnet masks, IP addresses and ports. To aggregate the similar alerts generated by the attacks which used the same router, host or port, the aggregation processes on the stages of subnet, host and service were respectively carried out based on attributes matching. On this basis, alerts correlation graph was generated by using one-step Markov chains model. In the graph, the directed edges represented the conditional transition probabilities between attack types, and the transition probabilities were calculated by the number of adjacent alerts. Finally, in the experiment, DARPA2000 traffic data was handled by the intrusion detection system Snort which was been configured as the most strict mode. After generating intrusion alerts set of LLDoS1.0 attack scenario, the above aggregation and correlation methods were conducted on the alerts of five types. The most ideal internal threshold of the self-extending time windows was further determined by parameter optimization. In this way, the alerts were reduced by the multistage aggregation effectively, and the results of aggregation were in accordance with the distribution of alerts source IP and source ports. Moreover, the accuracy rate of alerts correlation was calculated by comparing the correlation results with the official description of LLDoS1.0. Experiments demonstrated that the accuracy rate of the proposed method was 97.94%, which was 2.29% higher than that of traditional method.

Key words: network security; intrusion detection; alerts aggregation; time windows; multistage division

收稿日期:2016-09-15

基金项目:国家自然科学基金资助项目(61672531);湖北省自然科学基金资助项目(2015CFC867)

作者简介:李洪成(1991—),男,博士生。研究方向:网络安全;数据挖掘。E-mail:ytzyzbz@163.com

http://jsuese.ijournals.cn http://jsuese.scu.edu.cn

目前,随着网络攻击的隐蔽性和持续性越来越强,网络攻击检测工作面临的挑战日益严峻。入侵检测系统(intrusion detection system,IDS)通过对网络流量或数据包等数据的分析,可以有效地对网络攻击进行动态检测^[1-2]。但对于具有高隐蔽性和高持续性的复杂多步攻击,传统IDS普遍存在告警数据冗余度较大、数据可读性不强的问题^[3-5],严重影响安全管理员快速识别攻击行为和攻击意图。因此,有必要对相同或相似攻击行为所引发的重复告警进行有效聚合,为后续的网络告警关联分析和攻击场景识别提供基础^[6-7]。

国内外学者在网络告警聚合领域做了许多卓有成效的研究工作。文献[8]为了约简复杂的告警关联图,利用语义抽象的方法来把相同或相似的告警聚合,进而判断多个告警信息是否属于同一攻击步骤。但是该方法需要事先形成语义知识库,因此其可扩展性和实时性存在不足。

基于属性相似度的告警聚合方法可以有效弥补上述不足,该类方法首先由文献[9]提出,其基本思想是对告警数据的各维属性分为枚举型、连续型、布尔型等类别,分别根据属性的含义定义相似度计算函数,进而计算各告警的整体相似度,并聚合整体相似度超过阈值的告警。基于属性相似度的聚合方法存在的不足是告警中的时间、IP地址等属性在告警聚合中具有特殊作用,将其代入运算缺乏合理性^[10]。

针对此问题,很多学者对时间、IP地址等属性进行了单独研究。其中:在IP地址属性处理方面,文献[11]将源/目的IP相互关联的所有告警聚为一簇,在各簇内部进行告警聚合,但是在处理DDOS等完整攻击场景时,所有IP都将相互关联,导致告警被划分在同一簇中,进而使分簇失效,因此该方法的适用范围有限。在时间属性处理方面,相关研究方法大多集中在一定大小的滑动时间窗口内进行告警聚合,时间窗口的设定问题是其中的难点,窗口设置过大或过小都会影响告警聚合的有效性。如果时间窗口过大,则导致属于多个攻击行为的告警被聚合在一起,即增加错误聚合的个数;如果时间窗口过小,则无法聚合单次攻击行为所引发的所有告警^[12-14]。针对此问题,文献[15]考虑到持续性网络攻击的特点,对传统的滑动窗口方法进行改进,先利用较长的时间窗口存储当前的告警信息,再采用流式处理方法在滑动时间窗口内进行告警聚合。该方法的不足在于当前窗口和滑动窗口的大小仍需主

观确定。文献[16]为解决时间间隔波动较大的告警聚合问题,改进了基于固定时间阈值的聚合方法,每次产生新的告警时都会计算告警序列时间间隔的相对均方差,并以此作为告警聚合时间窗口的变异系数,可以实现时间阈值的动态更新。但是,该方法计算得出的动态时间窗口是单调递增的,只有当告警间隔的均方差趋于零时,时间窗口才会收敛,否则就会不断增大。该条件在实际中较难满足。

针对现有方法存在的不足,本文提出一种基于自扩展时间窗口的告警多级聚合与关联方法。在时间窗口设定方面,将时间间隔相近告警的时间窗口动态拉长;在IP地址属性处理方面,将端口、IP地址和所在子网相同的告警聚合在一起。进而,通过计算告警间的1阶状态转移概率进行告警关联,为攻击场景识别提供基础。

1 自扩展时间窗口设定方法

由于不同的入侵检测系统产生的告警格式不同,且在各项属性中很多属性对于告警聚合贡献不大,因此在进行告警聚合之前首先要规范告警的数据类型。本文使用现有最典型的告警消息格式IDMEF来描述入侵检测系统产生的网络告警,并在该数据结构下进行告警聚合。告警的数据结构为 $Structure = \{Time, SrcIP, DstIP, SrcPort, DstPort, AttackType\}$ 。其中, $Time$ 为告警所代表攻击的开始时间, $SrcIP$ 和 $DstIP$ 分别为攻击的源/目的IP, $SrcPort$ 和 $DstPort$ 分别为攻击的源/目的端口, $AttackType$ 为攻击类型。

在告警的各维属性中, $Time$ 属性对于告警聚合至关重要。对于同一攻击步骤所产生的重复告警,以及漏洞扫描、洪泛攻击等所产生的大量相似告警,其相邻告警两两之间的 $Time$ 属性间隔较小,而攻击的整体跨度呈现出不确定性^[17],因此,提出一种基于自扩展时间窗口的划分方法,利用告警的靠近关系进行时间窗口的自动扩展。其时间窗口设定的具体步骤如下:

Step 1:数据预处理。将告警开始时间 $Time$ 属性换算为以s为单位。

Step 2:将告警序列按 $Time$ 属性升序排列,并将第1个告警的开始时间设置为第1个时间窗口的起点。

Step 3:若后一个告警的 $Time$ 与前一个的 $Time$ 相差小于预先设定的时间间隔阈值 I ,则将后一个告警归入当前时间窗口中。否则,将前一个告警的 $Time$ 设置为当前时间窗口的终点,并将后一个告警的 $Time$

设置为下一个时间窗口的起点。具体算法如下。

算法 1 基于自扩展时间窗口的划分算法

输入: 告警集合 $Alert$, $Alert$ 中的告警数 N , 预设时间间隔阈值 I 。

输出: 经自扩展时间窗口划分后的告警集合, 存储于 3 维数组 $setB_j$, 其每一页为一个时间窗口内的告警。

```

1) Begin
2)  $p = 1$ ;
3)  $k = 1$ ;
4) for ( $i = 1; i \leq N; i++$ )
5)   if( $Alert_{i+1}.Time - Alert_i.Time \leq I$ ) then
6)      $setB_j(p, :, k) = Alert(i, :)$ ;
7)      $p++$ ;
8)   end if
9)   else then
10)     $p = 1$ ;
11)     $k++$ ;
12)     $setB_j(p, :, k) = Alert(i, :)$ ;
13)   end else
14) end for
15) return  $setB_j$ ;
16) End

```

2 网络告警多级聚合方法

在告警聚合过程中, 传统的基于属性相似度的方法把 $AttackType$ 、 $SrcIP$ 、 $DstIP$ 、 $SrcPort$ 和 $DstPort$ 的相似度作为整体相似度的一部分进行融合计算, 该方法存在一定局限性, 原因为: 1) $AttackType$ 区分了告警记录所属于的类别, 是告警信息记录中最重要的 1 维属性。将 $AttackType$ 代入相似度计算会将彼此无关的攻击产生的告警记录聚合在一起, 降低聚合的准确性。2) 往往把 IP 地址进行按位匹配, 只能一定程度上反映 IP 属于同一子网的概率, 而不能明确地判断其是否属于同一子网。3) 把 IP 和端口代入相似度计算的方法忽略了其本身代表的意义, 难以反映攻击的实际情况。

针对上述问题, 提出一种告警多级聚合方法。利用 $AttackType$ 对告警进行分簇, 避免不同 $AttackType$ 的告警相互聚合; 在 IP 地址和端口的处理方面, 利用子网掩码来明确地划分子网, 进而进行子网级、主机级和服务级聚合, 这 3 级聚合分别对应 3 种不同阶段攻击所产生的告警, 聚合的结果具有较强的实际意义。具体的聚合过程如下:

Step 1: 利用 $AttackType$ 进行第 1 层划分。将 $AttackType$ 相同的告警分在一个簇内, $AttackType_i$ 所对应的簇记为 $setA_i$ 。

Step 2: 基于自扩展时间窗口的告警划分。在 Step 1 划分出的簇 $setA_i$ 内, 利用算法 1 得到同一时间窗口内的告警集合, 第 j 个时间窗口内的告警集合记为 $setB_j$ 。

Step 3: 子网级聚合。在 Step 2 划分出的告警集合 $setB_j$ 内, 先利用子网掩码将 $SrcIP$ 属于同一子网的告警划分到一个集合中, 第 k 个子网内的告警集合记为 $setC_k$, 然后将该集合内的告警聚合起来。之所以利用 $SrcIP$ 所属的子网进行划分, 是为了反映攻击者利用同一个路由器作为跳板而实施的攻击。

Step 4: 主机级聚合。在 Step 3 划分出的告警集合 $setC_k$ 内, 将 $SrcIP$ 属性相等的告警划分到一个集合中, $SrcIP = IP_p$ 的告警集合记为 $setD_p$, 然后将该集合内的告警聚合起来。本聚合过程可以反映攻击者利用同一主机作为傀儡而实施的攻击。

Step 5: 服务级聚合。在 Step 4 划分出的告警集合 $setD_p$ 内, 将 $SrcPort$ 属性相等的告警划分到一个集合中, $SrcPort = Port_q$ 的告警集合记为 $setE_q$, 然后将该集合内的告警聚合起来。本聚合过程可以反映攻击者利用同一服务权限作为跳板而实施的攻击。

基于多级划分思想的网络告警聚合整体流程如图 1 所示。

在告警多级聚合流程中, 多次利用属性匹配进行告警划分, 具体算法如下。

算法 2 基于属性匹配的划分算法

输入: 待划分的一个告警集合 $Alert$, 属性可能取值的个数 H 。

输出: 根据属性划分后的告警集合, 存储于 3 维数组 X , 其每一页为一个集合内的告警。

```

1) Begin
2) for ( $h = 1; h \leq H; h++$ )
3)    $p = 1$ ;
4)   for ( $l = 1; l \leq N'; l++$ )
5)     if( $strcmp(Alert_l.Attribute', Attribute_h')$ ) then
6)        $X(p, :, h) = Alert(l, :)$ ;
7)        $p++$ ;
8)     end if
9)   end for
10) end for
11) return  $X$ ;
12) End

```

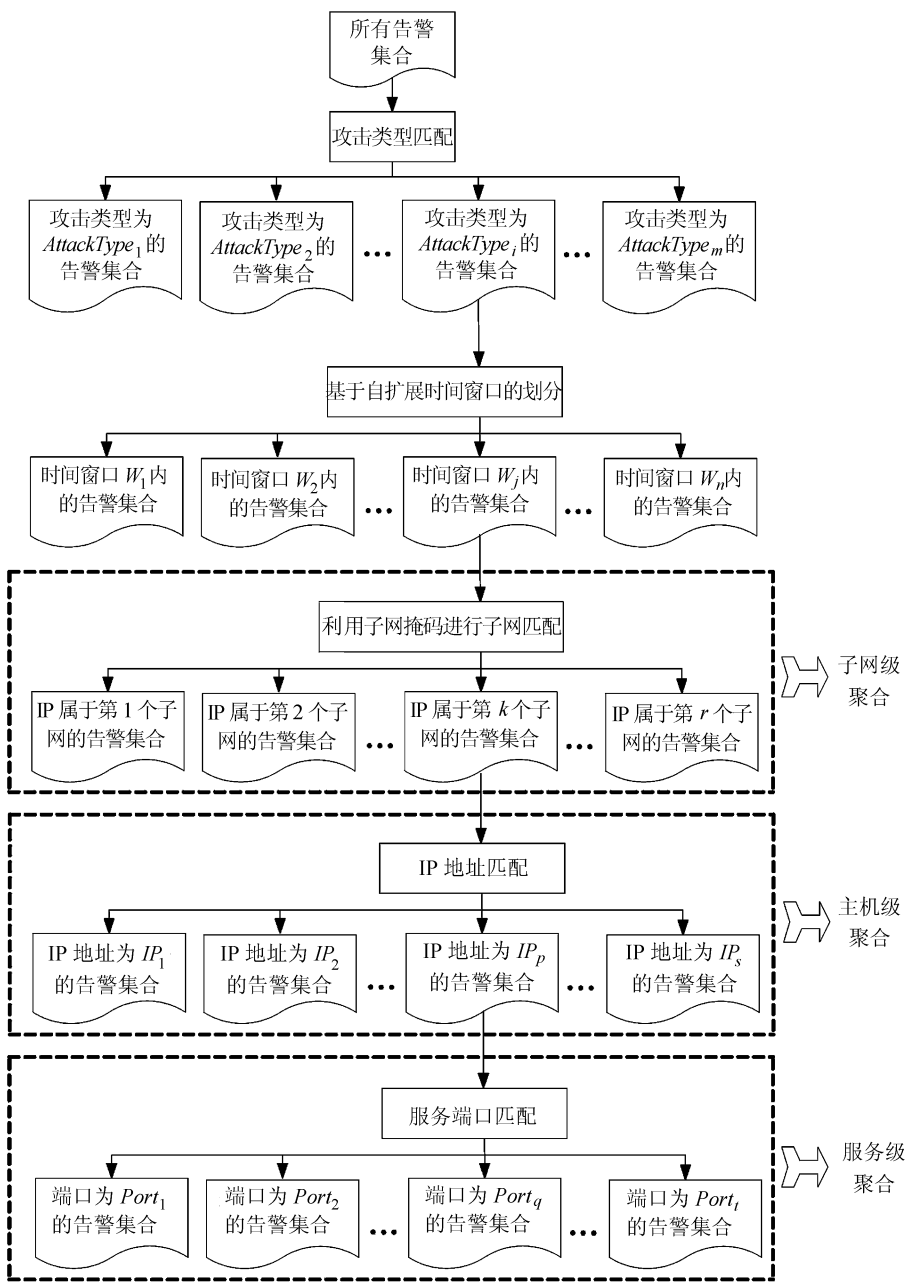


图 1 网络告警多级聚合整体流程图

Fig. 1 Integral flow chart of multistage aggregation for network alerts

3 网络告警关联图生成方法

攻击者实施的一个攻击场景可被看作多个逻辑步骤,当前的逻辑步骤(而不是时序步骤)使攻击者处于特定的权限状态,下一步可能采取的攻击行为只与当前所处的权限状态有关,而与之之前的权限状态无关,即攻击者下一步的攻击步骤与之前发生过的攻击路径无关。因此在逻辑步骤上,多步攻击满足 1 阶马尔可夫模型的无后效性。故可以利用 1 阶马尔可夫性质对多步告警关联过程进行建模。

定义(告警关联图) 生成的告警关联图是

一个边加权的有向图 $G(V,E)$ 。其中: V 为告警类型节点的集合,每个节点代表一种攻击类型,记为 $type_i$; E 是有向边 e_{ij} 的集合, e_{ij} 表示从节点 $type_i$ 指向 $type_j$ 的边; e_{ij} 的权重 p_{ij} 表示 $type_i$ 到 $type_j$ 的转移概率,即攻击者从 $type_i$ 转移到 $type_j$ 的条件概率, p_{ij} 的取值由 $type_i$ 和 $type_j$ 时序依次出现的概率决定。在告警关联图中,以同一节点为出发点的边权之和一定为 1。

告警关联图可以定量地反映攻击者的攻击路径和意图。衡量告警关联效果的指标主要为误关联率 FR 和漏关联率 MR , 其计算公式如下:

$$FR = \frac{N_{F_COR}}{N_{T_COR}}, MR = \frac{N_{M_COR}}{N_{T_COR}} \quad (1)$$

式中, N_{F_COR} 、 N_{M_COR} 和 N_{T_COR} 分别为误关联数、漏关联数和总关联数。此外, 关联准确性指标 *Accuracy* 可以对误关联率 *FR* 和漏关联率 *MR* 进行综合评价。*Accuracy* 取值越大, 表示告警关联的效果越好。*Accuracy* 的计算公式如下:

$$Accuracy = \frac{2 \times (1 - MR) \times (1 - FR)}{(1 - MR) + (1 - FR)} \quad (2)$$

4 实验及分析

在进行告警聚合之前需要首先利用入侵检测系统产生告警信息。本实验使用的网络流量数据为 DARPA 2000 数据集。DARPA 2000 是美国国防部对其内网主机模拟攻击的流量数据, 是进行入侵检测的常用数据集。本实验使用 DARPA2000 中 LL-DoS1.0 攻击场景的数据。该攻击场景是在多重网络和多重审计期中运行的一个 DDoS 攻击, 网络流量的格式为 tcpdump, 从林肯实验室的官方网站上

下载, 取其所有流量的前 10%。配置 Snort 为最严格模式, 加载 Snort 的规则表和匹配引擎, 处理. tcp-dump 文件, 可以输出告警数据共 4 399 条。

首先, 利用攻击类型对告警进行划分的结果如表 1 所示。

表 1 基于攻击类型的划分结果

Tab.1 Division results based on attack types		
告警的攻击类型名称	攻击类型	告警数量
Sadmind_ping	<i>AttackType</i> ₁	1 283
Sadmind_Amslverify_Overflow	<i>AttackType</i> ₂	276
Admind	<i>AttackType</i> ₃	87
Rsh	<i>AttackType</i> ₄	65
Mstream_Zombie	<i>AttackType</i> ₅	77

利用自扩展时间窗口方法对表 1 中各攻击类型的告警进行划分, 得到的结果如表 2 所示, 其中, n_{set} 为划分出的时间窗口个数, $n_{alert_{max}}$ 为时间窗口中最多包含的告警数。

表 2 基于自扩展时间窗口的划分结果

Tab.2 Division results based on self-expanding time windows										
攻击类型	<i>I</i> = 10 s		<i>I</i> = 30 s		<i>I</i> = 60 s		<i>I</i> = 90 s		<i>I</i> = 120 s	
	n_{set}	$n_{alert_{max}}$	n_{set}	$n_{alert_{max}}$	n_{set}	$n_{alert_{max}}$	n_{set}	$n_{alert_{max}}$	n_{set}	$n_{alert_{max}}$
<i>AttackType</i> ₁	101	1 151	92	1 151	80	1 151	67	1 152	58	1 165
<i>AttackType</i> ₂	213	2	176	5	142	8	103	9	86	10
<i>AttackType</i> ₃	50	7	20	11	14	25	12	25	9	42
<i>AttackType</i> ₄	33	5	27	5	25	6	23	6	22	6
<i>AttackType</i> ₅	64	2	56	3	46	4	40	4	35	6

为了对时间间隔阈值 *I* 进行参数寻优, 在 *I* 取不同数值时, 对告警进行 3 级聚合, 进而观察聚合效果的变化情况。告警聚合效果用聚合率 *Rate* 衡量, 聚合率的计算公式为:

$$Rate = \left(1 - \frac{N_{aggregation}}{N_{origin}}\right) \quad (3)$$

式中, $N_{aggregation}$ 为聚合产生的告警数, N_{origin} 为原始告警数。

设子网掩码为“255. 255. 255. 0”, 经过 3 级聚合后最终得到的告警聚合率如图 2 所示。

从图 2 可以看出, 随着时间间隔阈值 *I* 的增大, 最终聚合率持续上升, 这是因为 *I* 的增大会使自扩展时间窗口增大, 进而导致最终划分结果中每个时间窗中的告警增多, 所以聚合率增高。然而, 由于时间窗口的增大也会导致一部分无关的告警被聚合, 即会降低聚合准确性, 所以利用聚合率连线斜率明

显降低的拐点来确定 *I* 的取值。得到攻击类型 *AttackType*₁ ~ *AttackType*₅ 的告警聚合过程所选用的 *I* 值分别为 90、90、30、60、30 s。时间间隔阈值 *I* 优化后各攻击类型告警的 3 级聚合结果如表 3 所示。

表 3 时间间隔阈值 *I* 优化后的 3 级聚合结果

Tab.3 Three-stages aggregation results after the optimization of time internal threshold *I*

攻击类型	聚合产生的告警数		
	子网级	主机级	服务级
<i>AttackType</i> ₁	67	78	87
<i>AttackType</i> ₂	103	139	152
<i>AttackType</i> ₃	20	22	29
<i>AttackType</i> ₄	25	27	27
<i>AttackType</i> ₅	26	57	57

在得到精简告警的同时, 本实验得到的聚合结果还可以在一定程度上反映网络中 IP 扫描、端口扫描出现的可能性。即主机级聚合的聚合率越高, 表明网

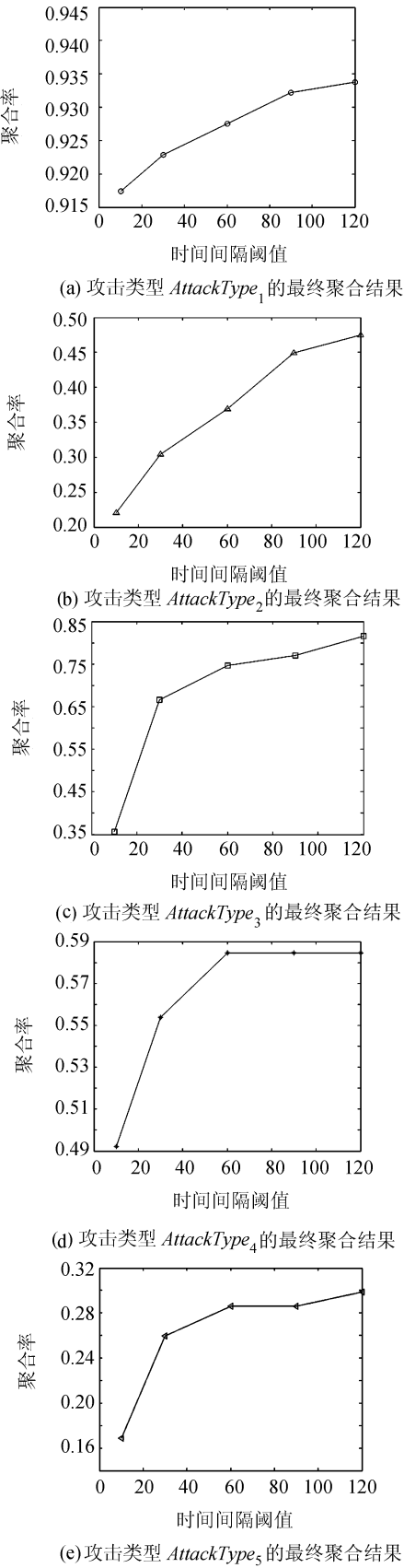


图 2 各攻击类型的最终聚合结果

Fig.2 Final aggregation results of attack types
络中出现 IP 扫描的可能性越高;端口级聚合的聚合

率越高,表明网络中出现端口扫描的可能性越高。

为了验证本文提出的告警聚合与关联方法的准确性,利用本文方法和文献[16]方法所得出的聚合结果分别进行告警关联。利用本文方法聚合时采用 I 值优化后的服务级聚合结果。采用文献[16]聚合方法和本文聚合方法得出的 LLDoS1.0 攻击场景的告警关联图分别如图 3 和 4 所示。

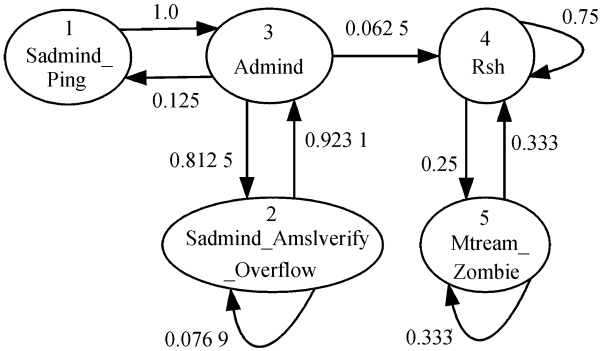


图 3 利用文献[16]聚合方法得到的关联图
Fig.3 Correlation graph obtained by using the aggregation method of ref. [16]

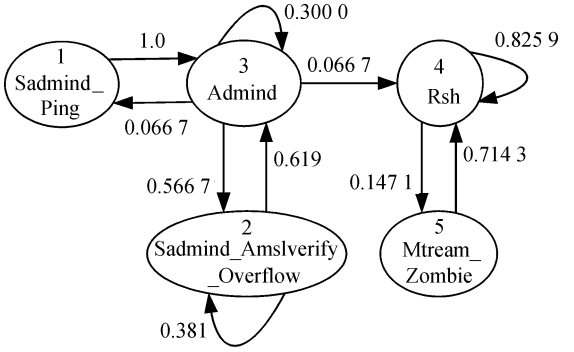


图 4 利用本文聚合方法得到的关联图
Fig.4 Correlation graph obtained by using the proposed aggregation method

在图 3、4 中,攻击类型“Sadmind_Amslverify_Overflow”和“Admind”之间以及“Rsh”和“Mstream_Zombie”之间均存在相互指向的关系,说明两类攻击在一段时间内反复出现。由于它们属于同一攻击阶段中的攻击行为,所以攻击者在同一内网中会交替尝试这些行为。另外,攻击类型由“Admind”到“Sadmind_Ping”的关联与官方公布的 LLDoS1.0 攻击场景不相符,故将其判定为误关联。通过计算得到,图 3 所示关联图与图 4 所示关联图对应的准确率 $Accuracy$ 分别为 95.65% 和 97.94%,本文方法的关联准确性高于文献[16]算法。这是因为本文的告警聚合方法粒度更细,时间窗口的设定更加合理,因此更能对相似告警进行有效聚合,减少了误报在

告警中所占的比例。

5 结 论

为提升攻击场景识别的可读性,本文在时间窗口的设置方面采用自扩展的策略,有利于聚合有计划性的持续攻击产生的告警;利用源/目的 IP 和源/目的端口等信息对告警进行多级聚合,以满足不同层次告警聚合的需要;在聚合告警的基础上进行了多步告警关联分析。实验证明本文方法具有良好的告警聚合效果和较高的关联准确性。下一步的研究工作主要是如何将聚合和关联方法扩展到更大规模的告警数据集上,并提高处理的时效性。

参考文献:

- [1] Kaynar K, Sivrikaya F. Distributed attack graph generation [J]. IEEE Transactions on Dependable and Secure Computing, 2015, 13(5): 519–532.
- [2] Zhang Yu, Liu Qingzhong, Li Tao, et al. Danger theory based real-time response model for APT attacks [J]. Journal of Sichuan University: Engineering Science Edition, 2015, 47(4): 83–90. [张瑜, Liu Qingzhong, 李涛, 等. 基于危险理论的 APT 攻击实时响应模型 [J]. 四川大学学报(工程科学版), 2015, 47(4): 83–90.]
- [3] Ma Dong, Wang Yongjun, Fu Zhenlong. A synergetic pattern matching method based-on DHT structure for intrusion detection in large-scale network [J]. Procedia Engineering, 2011, 15: 3511–3515.
- [4] Ramaki A A, Amini M, Atania R E. RTECA: Real time episode correlation algorithm for multi-step attack scenarios detection [J]. Computers & Security, 2015, 49: 206–219.
- [5] Zheng Ya, Chen Xingshu, Yin Xueyuan. Distribution denial of service detection algorithm based on PCC time series analysis [J]. Journal of Sichuan University (Engineering Science Edition), 2015, 47(Supp 2): 142–148. [郑亚, 陈兴蜀, 尹学渊. 基于 PCC 时间序列的 DDoS 检测算法 [J]. 四川大学学报(工程科学版), 2015, 47(增刊 2): 142–148.]
- [6] Li Longying, Li Jinku, Ma Jianfeng, et al. Comprehensive analysis of real-time alerts with attack strategy graph [J]. Journal of Xidian University, 2014, 41(5): 84–90. [李龙营, 李金库, 马建峰, 等. 采用攻击策略图的实时警报综合分析方法 [J]. 西安电子科技大学学报(自然科学版), 2014, 41(5): 84–90.]
- [7] Shittu R, Healing A, Ghanea-Hercock R, et al. Intrusion alert prioritisation and attack detection using post-correlation analysis [J]. Computers & Security, 2015, 50: 1–15.

- [8] Ning Ping, Cui Yun, Reeves D S, et al. Techniques and tools for analyzing intrusion alerts [J]. ACM Transactions on Information and System Security (TISSEC), 2004, 7(2): 274–318.
- [9] Valdes A, Skinner K. Probabilistic alert correlation [C]// Proceedings of 4th International Symposium on Recent Advances in Intrusion Detection (RAID 2001). London: Springer-Verlag, 2001: 54–68.
- [10] Elshoush H T I. An innovative framework for collaborative intrusion alert correlation [C]// Proceedings of 13th International Symposium on Science and Information. London: IEEE, 2014: 607–614.
- [11] Feng Xuewei, Wang Dongxia, Huang Minhuan, et al. A mining approach for causal knowledge in alert correlating based on the Markov property [J]. Journal of Computer Research and development, 2014, 51(11): 2493–2504. [冯学伟, 王东霞, 黄敏恒, 等. 一种基于马尔可夫性质的因果知识挖掘方法 [J]. 计算机研究与发展, 2014, 51(11): 2493–2504.]
- [12] Dong Xiaomei, Zhao Qian, Li Xiaohua, et al. Network forensics based on scenario reconstruction and alert aggregation [J]. Control and Decision, 2014, 29(1): 39–44. [董晓梅, 赵茜, 李晓华, 等. 基于场景重构与报警聚合的网络取证分析技术 [J]. 控制与决策, 2014, 29(1): 39–44.]
- [13] Ghasemigol M, Ghaemi-Bafghi A. E-correlator: An entropy-based alert correlation system [J]. Security and Communication Networks, 2015, 8(5): 822–836.
- [14] Fredj O B. A realistic graph-based alert correlation system [J]. Security and Communication Networks, 2015, 8(15): 2477–2493.
- [15] Qiu Hui, Wang Kun, Yang Haopu. Network alerts depth information fusion method based on time confrontation [J]. Journal of Computer Applications, 2016, 36(2): 499–504. [邱辉, 王坤, 杨豪璞. 基于时间对抗的网络报警深度信息融合方法 [J]. 计算机应用, 2016, 36(2): 499–504.]
- [16] Yan Shaohua. Research on alert fusion technology of network intrusion detection system [D]. Shenyang: Shenyang Aerospace University, 2011: 23–28. [晏少华. 网络入侵检测系统中报警数据融合技术研究 [D]. 沈阳: 沈阳航空航天大学, 2011: 23–28.]
- [17] Bateni M, Baraani A. Time window management for alert correlation using context information and classification [J]. International Journal of Computer Network and Information Security, 2013, 5(11): 9–16.

(编辑 赵 婧)

引用格式: Li Hongcheng, Wu Xiaoping. Multistage aggregation and correlation for network alerts based on self-extending time windows [J]. Advanced Engineering Sciences, 2017, 49(1): 206–212. [李洪成, 吴晓平. 基于自扩展时间窗的告警多级聚合与关联方法 [J]. 工程科学与技术, 2017, 49(1): 206–212.]