

一种基于模糊综合评判的入侵异常检测方法¹

张剑, 龚俭

(东南大学计算机科学与工程系, 南京 210096)

(zhangj@njnet.edu.cn; jgong@njnet.edu.cn)

摘要 目前国际上已实现的大多数入侵检测系统是基于滥用检测技术的, 异常检测技术还不太成熟, 尤其是基于网络的异常检测技术, 如何提高其准确性、效率和可用性是研究的难点。本文提出一种面向网络的异常检测算法 FJADA, 该算法借鉴了模糊数学的理论, 应用模糊综合评判工具来评价网络连接的“异常度”, 从而确定该连接是否“入侵”行为。实验证明, 该方法能检测出未知的入侵方式, 而且准确性较高。

关键词 入侵检测; 滥用检测; 异常检测; 模糊综合评判; 语言变量

中国法分类号 TP393.08

An Anomaly Detection Method based on Fuzzy Judgement

ZHANG Jian, and GONG Jian

(Department of Computer Science and Technology, Southeast University, Nanjing 210096)

Abstract Currently most implemented IDSs in the internation are based on the technology of Misuse Detection. The technology of Anomaly Detection, especially which is based on network, is not so mature. How to improve its accuracy, efficiency, and usability is the difficulty of research. In the paper, we present an algorithm of Anomaly Detection towards computer network: FJADA, which has benefited from the Fuzzy Mathematics. The algorithm applies Fuzzy Judgement to evaluate the Anomaly Degree of an Network Connection, then decides the Network Connection whether intrusivly or not. Our experiment has verified that the method can detect unknown intrusion and the accuracy is high.

Key words intrusion detection, misuse detection, anomaly detection, fuzzy judgement, linguistic variable

一. 引言

随着计算机网络的不断发展, 其自身的安全形势愈来愈严峻, 针对计算机网络的攻击事件层出不穷, 导致网络瘫痪和资源被盗用的事件给政府和企业带来了不可估量的损失。入侵是指破坏资源保密性、完整性和可用性的一组行为, 入侵检测则是指发现或检测入侵。目前针对它的研究是一个热点。

入侵检测方法分为两大类, 一种通过对采集的信息按已知的知识进行分析, 发现正在发生和已经发生的入侵行为, 这种方法被称为滥用检测方法; 另外一种方法是通过采集和统计发现网络或系统中的异常行为, 然后向管理员提出警告, 这种被称为异常检测方法。异常检测方法能检测出未知的入侵方式, 但其误报率较高, 如何降低该方法的漏检率 (false negative rate) 和误报率 (false positive rate) 是一个难点。

对异常检测方法的研究已经有很多。SRI 的 IDIES 和 NIDES 系统通过统计方法建立用户的正常行为映像, 并用它来比较用户的当前行为, 当发生偏差时, 认为有异常情况发生, 产生报警^[1]。Wenke Lee^[2]等人在 1998 年和 1999 提出利用数据挖掘技术来分析网络报文的特

¹ 本文受国家自然科学基金项目资助(90104031)

征并建立入侵检测模型。Lane 和 Brodley 在 1997 年提出用机器学习方法建立异常检测模型^[3]。上述的方法都存在一个共同的缺点：对建立检测模型的训练数据要求比较高，一方面要求它们是干净的数据（不包含入侵行为），另一方面又要求它们包含检测对象的大多数正常行为，而通常这两者是不可兼得的。而通过数据挖掘和机器学习方法建立的检测模型存在误报率比较高的缺点，因为它们将检测对象行为强行分为正常和入侵两类，或者将所有的异常行为都划入入侵行为中。

将模糊数学的理论和方法运用到异常检测技术是最近两年兴起的，Dickerson 等人在 2000 年和 2001 年提出一个基于代理的异常检测系统：模糊入侵识别引擎（FIRE）^[4-5]，FIRE 用数据挖掘技术处理网络数据，从中提取出对异常检测有意义的测度，然后建立这些测度的模糊集合。FIRE 用一个模糊分析引擎评价这些模糊集合，并据此产生响应。Siraj 等人在 2001 年提出建立模糊认识图（Fuzzy cognitive map）和模糊规则库来支持因果式知识推理过程^[6]；李之棠等人在 2000 年提出一个模糊入侵检测模型^[7]。但这些研究大多没有提出具体的检测算法和实验结论。

本文依托 CERNET 的网络环境，在模糊数学的理论基础上提出了一种具体的异常检测算法：基于模糊综合评判的异常检测方法 FJADA(Fuzzy Judgement-based Anomaly Detection Algorithm)，理论和实验表明，该算法能检测到面向网络连接的未知攻击方式。该算法的检测模型仅要求训练数据是真实且足够多；另外，考虑到网络中存在大量介于正常和入侵之间的异常行为，本文采用了基于模糊综合评判的模糊分析引擎来减少误报率。

本文的组织结构如下：第二节介绍 FJADA 算法的理论依据和设计思想；第三节介绍算法的具体实现；第四节对算法的实现方法进行局部的细节讨论；实验结果将在第五节给出；最后是总结。

二. FJADA 算法的理论依据和设计思想

FJADA 以网络连接为检测对象，运行于 Internet 与内部网间的网关或与之相联的内部计算机中，目标是检测基于网络的未知攻击方式。本节首先介绍必要的模糊数学概念和应用，然后介绍 FJADA 的检测模型，最后是 FJADA 的设计思想。

2.1 模糊数学相关概念和应用

2.1.1 模糊数学基本概念^[8]

(1) 模糊集合和隶属度

模糊集合用来描述一个模糊概念，它是内涵和外延都不明确的集合。一个模糊集合 A 可用隶属度函数 μ 来表示， μ 的定义域是与该模糊概念相关的对象的集合，称为 A 的论域 U ，而 μ 表示了 U 中每个对象隶属于 A 的程度，称为隶属度， $\mu \in [0, 1]$ 。模糊集合是模糊数学的基本概念，通过它可以建立与普通集合的关系，继而用经典数学的方法来研究模糊现象。

(2) 水平集合（ α -level）

模糊集合 A 的一个 α 水平集合是一个普通集合记为 A_α ，它由 U 上所有对 A 的隶属度大于或等于 α 的元素组成。

(3) 模糊关系

模糊关系是关系的推广，在模糊关系中的两组事物的关系不是仅用“有”或“无”来描述，而是用 0 到 1 之间的数代表隶属于该模糊关系的程度。

(4) 语言变量

语言变量与数值变量不同，语言变量的值不是数而是自然语言或人工语言的字或句，称为语言值。语言值的语义由模糊集合来描述，因此，语言变量取模糊集合作为它的值。

2.1.2 模糊综合评判

模糊综合评判是模糊数学的重要应用,其目的是解决所研究的课题在影响因素十分复杂的情况下,选择最优方案或是按评价结论对系列产品质量(或人的、国家的综合能力)进行排序而决定取舍。

一般来说,模糊评判方法要有如下主要步骤:(1)确定评价因素集,评语集;(2)建立各因素的模糊集合(隶属函数);(3)建立评价因素与评语之间的模糊关系;(4)确定各因素在评判中所占的权重;(5)按一定的算子进行计算得到评判结论。

2.2 FJADA 的检测模型

定义 1 设网络连接特征的取值范围是 U , 对于 $\forall x \in U$, 其密度分布函数 $f(x)$ 是特征值为 x 出现在所有网络连接中的比例。

比较特殊的是建立量化尺度的特征的密度分布函数,在实际应用中一般先对其值域划分区间,然后计算区间的密度分布函数。这主要因为我们只关心在特征值在一定范围的出现比例。

FJADA 的检测模型是网络连接各种特征的密度分布函数的集合。它的特点是:

- (1) 实现简单,训练周期短,不要求训练数据是干净的。后者是因为某个测度值或范围的出现频率越高,则它的异常度就越低;而频率越低,则其异常度就越高,即越值得怀疑。这个思想是基于这样一个统计规律:在真实的网络环境中,从网络流量的长期¹变化趋势来看,发生在其中的入侵行为是占少部分的,大部分还是正常的行为。
- (2) 该模型是比较稳定的。只要训练数据量足够,该检测模型是比较稳定和精确的,因为它反映了特定网络的宏观行为规律,这在短期内是不会有很大变化的。
- (3) 根据网络连接特征的密度分布函数建立测度集与“异常度”评语集的模糊关系矩阵,然后用模糊综合评判的方法来评估整个连接的异常度属于哪个评语,并据此来确定该连接是否“入侵”行为。这种方法在实践中证明能有效地降低误报率。

2.3 FJADA 算法设计思想

异常检测中的“异常”实质是一个模糊概念,因为检测对象的行为是如此复杂,以致于无法给“异常”下准确的定义。异常检测的目的是发现异常的行为,然后根据异常的程度来决定该行为是否是入侵行为。这种思路用模糊数学来描述非常地适合,例如在基于网络的入侵检测中,可以用模糊关系来描述网络连接的异常度,而入侵行为则是异常度在某个评语以上的网络连接的集合。这说明将异常看成是一个模糊概念,用模糊关系来描述检测对象的异常度并用模糊数学的方法来实现异常检测算法是理论可行的。FJADA 算法以网络连接作为检测对象,用模糊关系来描述网络连接的异常度,在计算出网络连接的各种特征值对不同异常度评语的隶属度后,运用模糊综合评判计算出该连接的异常等级,从而决定它是否入侵行为。

FJADA 算法借鉴了模糊数学中的许多概念。除了模糊集合、模糊关系等基本概念,还用到语言变量。将网络连接的“异常度”看成是语言变量,其语言值“正常”、“有点异常”和“异常”等是评价网络连接对象异常的程度或等级,那么利用它们就能构成模糊综合评判的评语集,如果某个网络连接与某个评语的模糊关系隶属度比其他的评语都高,该网络连接的“异常度”就被评定为该评语,这称为最优隶属原则;而该网络连接的异常程度到达“异常”以上就可确定它为入侵行为。

本算法还利用模糊综合评判方法作为 FJADA 算法的基本框架。模糊综合评判的作用是能将对一件比较复杂的事件或事物的整体评估分成许多比较简单的小部分来分别评估,最终得出对整个事件或事物的结论。由于无法直接评估一个网络连接的“异常度”,就先分别评估它的每个测度的“异常度”,并据此建立模糊评判矩阵,然后用模糊综合评判得出该连接的“异常度”。本算法的建立检测模型阶段首先确定网络连接的特征集

¹ 按照网络行为学的定义,长期是指 1 天以上,而短期是指 10 分钟以下

合、特征权重向量和评语集合，然后在训练数据中建立每种特征的密度分布函数和据此建立与每个评语的模糊关系；在检测阶段，首先计算当前网络连接每个特征对各个评语的模糊关系隶属度，组成模糊关系矩阵，然后进行模糊关系合成运算得出当前连接的评语向量，最后根据最优隶属原则确定其“异常度”等级和确定它是否入侵行为。

三. FJADA 算法的具体实现

在前面已经介绍了模糊综合评判的步骤和作用，FJADA 算法基本遵循了它的整个过程，下面将按照模糊综合评判的步骤介绍。

3.1 确定评价因素集、评语集

FJADA 算法的第一步是确定评价因素集和评语集。网络连接可以用多个特征来描述，例如：连接持续时间、报文数量、平均报文长度、源地址、目标地址，服务类型（即目标端口号）等，这些特征的集合被称为评价因素集。评价因素分为两种，一些因素是否异常是与它的密度分布函数相关的，例如服务类型，如果它的值是 **80**，那么单纯从该测度来看，该连接是比较正常的，因为针对 **80** 端口的访问在网络流量中占的比例很大。这种因素被称为异常一致性因素。而有一些因素是否异常与其概率密度函数无关或部分无关，例如源地址、目标地址，这种因素被称为非异常一致性因素。对这种因素与“异常度”评语的模糊关系的建立在本文的第四部分有讨论。

评价向量是二元组 $E = \langle U, W \rangle$ ，其中 U 是评价因素集： $U = \{u_1, u_2, \dots, u_n\}$ ；而 W 是权重向量，其每个分量对应于一种因素在评判中的重要性比重，可表示为：

$$W = \int_U w/u$$

对应于评价因素集的是评语集，这里的评语集是“异常度”这个语言变量的若干语言值的集合。以后将用 $E = \{e_1, e_2, \dots, e_m\}$ 表示。

3.2 建立评价因素的“模糊集合”（“隶属函数”）

网络连接的评价因素集中的元素本质上不是模糊概念而是随机量，不应该用模糊集合来描述它们，但模糊现象与随机现象都是属于不确定的事件，它们的描述方式是一样的。因此，可以把这些因素的密度分布函数做为这些因素的“隶属函数”，这样这一步的工作就是从训练数据中计算每个因素的密度分布函数。对于非异常一致性因素无需这样做，因为不需要用它们的密度分布函数来确定与“异常度”评语的模糊关系。

3.3 建立评价因素与评语之间的模糊关系

评价因素集合是网络连接特征的集合，而评语则是对这些特征值异常程度的自然语言描述，例如“正常”、“异常”、“非常异常”等等。一般对于某个特征值容易判断它是否异常，而很难说它是属于哪个异常程度的，而用“属于某种异常程度的可能性”就比较适合描述这种情况，因此可用模糊关系来表示每个特征对每种评语的隶属度。

这是最重要的一步，一旦两者的模糊关系确定了，检测模型就基本上建立起来了。评价因素与评语之间的模糊关系表示各评价因素属于各异常程度的隶属度确定评价因素 u_i 与评语 e_j 之间的模糊关系是以 u_i 的密度分布函数 $f(u_i)$ 为根据的，假如评语集是 $\{e_1, e_2, \dots, e_m\}$ ，那么 u_i 的密度分布函数就分别映射到 m 个模糊关系，这些模糊关系隶属函数之间的关系如图 1 所示：

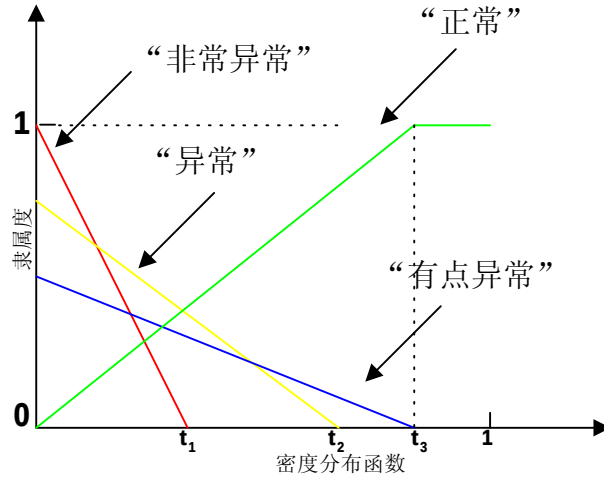


图1 各种评语的隶属函数

定义 2 若某网络连接特征 u_i 的密度分布函数为 $f(x)$, $x \in U$, U 是 u_i 的取值范围, 而 $f(x)$ 到某评语 e_j 的映射为 g , 则在 u_i 的值域下该评语的隶属函数 $\mu_{e_j}(x) = g(f(x))$, $x \in U$, 这也就是它们之间的模糊关系。

由图 1 可以得到评价因素与各个评语的模糊关系确定的两点原则:

原则 1. 某特征值的密度越小, 它对表示异常程度高的评语的隶属度越高。

当密度大于阈值 t_3 时, 该特征值被认为不异常, 因此与相对应的“异常度”评语的模糊关系隶属度为 0, 与“正常”的模糊关系隶属度为 1。当密度在区间 $[0, t_3]$ 中时, 特征值与“有点异常”的隶属度 > 0 , 但随着密度的降低隶属度上升很缓慢, 当密度降低到一定程度 ($[t_1, t_2]$ 间的某个值) 后, 它对“异常”的隶属度大于“有点异常”。即随着密度的降低到一定程度 (“异常”曲线与“有点异常”曲线的交点), 它属于异常程度高的可能性就较高。

原则 2. 表示异常程度越高的评语的隶属函数斜率越高

例如对于“非常异常”这个评语, 其隶属函数只在 $[0, t_1]$ 区间大于 0, 但随着密度的降低, 其隶属度很快超过“异常”。

图 1 中表现的评价因素与各个评语模糊关系隶属函数之间的关系是符合逻辑和实际的, 它是网络连接特征异常性的客观反映。

图 1 是用线性函数来描述概率密度函数与评语的模糊关系的示例, 而在实际中, 应该根据实验效果来调整模糊关系的数学模型, 但是需要满足上述两条原则。

3.4 确定各因素在评判中所占的权重

3.4.1 最大特征根法

确定各因素在评判中所占的权重是比较困难的, 这是因为

- (1) 对于检测不同的攻击类型, 特征权重向量可能不同。

定义 3 设网络连接特征集合 $U = \{u_1, u_2, \dots, u_n\}$, 则其特征权重向量为

$W = [w_1, w_2, \dots, w_n]$, 其中 w_i 为特征 u_i 的权重。

例如在检测拒绝服务攻击中, 报文数和报文长度占的权重就较大; 而在检测基于 Web 的攻击时源/宿 IP 地址和端口所占的权重就较大。而对于未知的攻击类型各特征所占的比重实际是不可知的。

解决这个问题的方法是把所有攻击方式按不同性质分类,例如在网络入侵检测系统重可将攻击方式分为拒绝服务攻击 (Dos)、缓冲区溢出攻击 (Buffer Overflow Attacks)、扫描攻击 (Scan)、解码类攻击 (Decoding)、基于 Web 的攻击和其他类型攻击。每类攻击的特征权重向量是唯一的,这样只需求出每种攻击类型的特征权重向量。

(2) 确定每个因素在整体中的顺序是困难的

确定各因素的权重实际是对所有因素的重要性进行排序,而在因素比较多时这是比较困难的,而如果只对任两个因素进行比较则是容易的,本文引入矩阵代数中的最大特征根法,该方法从任两个因素之间的相对权重组成的矩阵出发,求出各因素在总体中的权重。下面是该方法的描述:

最大特征根法首先确定一个比较判断矩阵,比较判断矩阵的内容是各因素之间相对权重,通过计算该矩阵的最大特征根所对应的特征向量,这个特征向量就是特征权重向量。假定评价因素集中有 n 个因素,它们的重要性权值分别是 $W_1, W_2, \dots, W_n$, 他们满足 $W_1 + W_2 + \dots + W_n = 1$ 。比较它们之间的重要性,很容易得到它们之间逐对比较的判断矩阵:

$$A = \begin{bmatrix} \frac{W_1}{W_1} & \frac{W_1}{W_2} & \dots & \frac{W_1}{W_n} \\ \frac{W_2}{W_1} & \frac{W_2}{W_2} & \dots & \frac{W_2}{W_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{W_n}{W_1} & \frac{W_n}{W_2} & \dots & \frac{W_n}{W_n} \end{bmatrix} = (a_{ij})_{n \times n}$$

$$\begin{aligned} \text{显然} \quad a_{ij} &= 1/a_{ji} \\ a_{ii} &= 1 \\ a_{ij} &= W_i / W_j \quad i, j = 1, 2, \dots, n \end{aligned}$$

定理 1 设特征权重向量 $W = [W_1, W_2, \dots, W_n]^T$, 则它是比较判断矩阵的最大特征向量

证明: 用 W 右乘 A 矩阵, 其结果为:

$$AW = \begin{bmatrix} \frac{W_1}{W_1} & \frac{W_1}{W_2} & \dots & \frac{W_1}{W_n} \\ \frac{W_2}{W_1} & \frac{W_2}{W_2} & \dots & \frac{W_2}{W_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{W_n}{W_1} & \frac{W_n}{W_2} & \dots & \frac{W_n}{W_n} \end{bmatrix} \begin{bmatrix} W_1 \\ W_2 \\ \vdots \\ W_n \end{bmatrix} = \begin{bmatrix} nW_1 \\ nW_2 \\ \vdots \\ nW_n \end{bmatrix} = nW$$

从上式不难看出, 特征权重向量 W 是比较判断矩阵 A 的对应于 n 的特征向量。根据矩阵理论可知, n 为上述矩阵 A 唯一非零的, 也是最大特征根, 而 W 则为其所对应的特征向量。命题得证。

因此问题就转化为求比较判断矩阵的最大特征向量, 使用矩阵代数的算法容易实现该目标。

3.4.2 建立比较判断矩阵

本文采用专家评判法来建立比较判断矩阵，专家评判法是模糊数学中建立模糊集合、模糊关系等数学模型的重要手段，它主要依赖于相关领域的专家的经验。用专家评判法建立比较判断矩阵的步骤如下：

(1) 请 n 位专家针对某种攻击类型根据自己的经验建立比较判断矩阵 $A_1, \dots, A_n$ 。

(2) 根据专家的权威性高低设置一组权值 $W_1, \dots, W_n$, $W_1 + \dots + W_n = 1$, W_i 代表第 i 位专家的权威性, $i \in 1, \dots, n$ 。

(3) 最终的比较判断矩阵 $A = W_1 * A_1 + \dots + W_n * A_n$ 。

3.5 按一定的算子进行计算得到评判结论

评判计算的步骤是：

(1) 根据 3.3 节确定的模糊关系对每个特征值根据评语集进行评定，最终合成评定矩阵。

(2) 用评价因素集的权重向量与评定矩阵进行模糊矩阵的合成运算，得出的向量称为综合评判向量。

(3) 根据最大隶属度原则，确定该特征集的评语。

设：特征集为 F , $F = \{f_1, f_2, \dots, f_n\}$ ；评语集 $E = \{e_1, e_2, \dots, e_m\}$ ，对单因素进行评定，得到评定向量 $R_i = (r_{ij})$, $i=1, 2, \dots, n$, $j=1, 2, \dots, m$ 。则评定矩阵

$R^T = [R_1 \ R_2 \ \dots \ R_n]$ 。设特征权重向量 $W_F = (W_{f1}, W_{f2}, \dots, W_{fn})$ ，这样即可按模糊

合成，求得模糊综合评判矩阵

$$S = (S_1, S_2, \dots, S_n)$$

$$= (W_{f1}, W_{f2}, \dots, W_{fn}) \circ \begin{bmatrix} r_{11} & r_{12} & \dots & r_{1m} \\ r_{21} & r_{22} & \dots & r_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ r_{n1} & r_{n2} & \dots & r_{nm} \end{bmatrix}, \text{ 设 } S_{k0} = \max(S_i), i=1, 2, \dots, n. \text{ 则该}$$

评价因素集属于 k_0 等级。

3.6 确定入侵事件

入侵事件 I 是评语集的一个子集，一般将评语为“异常”程度以上的网络连接划分为入侵事件。

四。算法局部讨论

4.1 非异常一致性因素与“异常”模糊关系的确定

在评价因素集中，存在一些非异常一致性因素，例如源和宿 IP 地址、源和宿端口号等。对于这种因素，一种简单的处理办法就是设它们与所有评语之间的模糊关系的隶属度为 0。但在实际网络环境中容易发现：某台主机发起的网络连接是入侵事件，那么该主机很可能就是攻击的发起者，其地址的“异常度”就会增加。基于这个发现，可以认为非异常一致性因素的“异常度”与发现到的入侵事件数量有关，即：

$$m_{\text{“异常”}}(x) = f(\text{EventNum}), x \in U, 0 \leq f(\text{EventNum}, t) \leq 1, \text{ 其中 } \text{EventNum} \text{ 是检测到}$$

的入侵事件数量。设入侵事件被描述为： $E = (\text{事件性质 } I, \text{发现时间 } T, \text{非异常一致性因素集 } F)$ 。设相同非一致性因素的入侵事件序列为 $ES = \{e_1, e_2, \dots, e_n\}$, $|e_i.T - e_{i+1}.T| \leq \alpha$, $i=1, 2, \dots, n-1$ 。则 $\text{EventNum} = |ES|$ 。

4.2 相互依赖因素与评语集模糊关系的确定

有些异常一致性因素是否异常不是只取决于它的概率密度函数，而是与其他因素有关。例如对于 *ftp* 连接，它的平均报文长度就要比 *Telnet* 连接的平均报文长度要大，如果不管是

什么应用,单纯根据概率密度函数来决定平均报文长度的异常度,那么检测的精确性就较低。解决这个问题的方法是将相互依赖的因素联合作为一个复合因素来计算其联合概率密度函数,然后再按第3节的方法来确定它与评语之间的模糊关系。

五. 实验及其结论

实验过程分为获取网络数据、对数据进行预处理、建立检测模型和入侵检测。实验框架如图2所示:

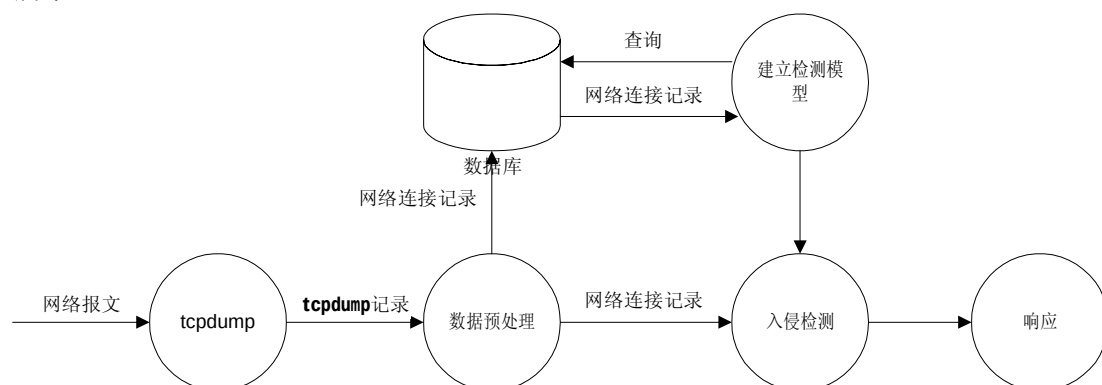


图2 入侵检测系统框架

实验结果将与 Lee 在 1999 年的实验结果相比较,因此需要简单地回顾一下 Lee 的实验过程。

5.1 Lee 的异常检测实验

Wenke Lee 在 1999 用数据挖掘技术建立入侵检测的分类模型,该模型以网络连接作为检测对象,首先将用 tcpdump 收集的四组网络数据整理成网络连接记录,一组被标识为“正常”的 80%数据用来建立分类模型,其余 20%和另外三组都被注入了不同类型的入侵事件的数据被用来测试检测效果。分类模型以网络连接的“服务类型”作为分类标准建立好分类模型后,在检测期将网络连接通过该模型,一旦发现结果与实际的服务类型不符合,则判断为入侵行为,即将分类错误率作为检测率或误报率。

5.2 数据预处理

本实验第一步是启动网关内的 tcpdump 进程收集网络数据,该进程的输出是多条网络报文的记录,每条记录是一个报文的测度多元组: <时间戳、源地址、源端口、目标地址、目标端口、连接标识、报文序列号、期望应答序列号、接收窗口大小、urgent 标识>。数据预处理的结果是生成网络连接的测度多元组: <连接起始时间、连接时间、源地址、目标地址、源-宿发送字节总数、宿-源发送字节总数、服务类型(即目标端口号)、协议执行状态>,该测度多元组与 Lee 的大致相同(多了“宿-源发送字节总数”),目的是以后可以排除其他因素而比较两者方法的优劣。同样是为了比较的需要,不从网络中直接获取网络数据,而是使用 Lee 的实验数据¹,该数据分为四组,第一组叫 baseline,是不包含模拟入侵活动的网络数据,以其中 80%作为训练数据,其余 20%用做测试误报率;第二、三和四组数据分别叫 network1、network2 和 network3,分别注入了三种不同类型的攻击活动,用做测试检测率。四组数据预处理后存入数据库的表的名称分别为 normal, intrusion1, intrusion2 和 intrusion3。

模仿 Lee 的实验将网络连接分为三类,一类是本地网络外出的网络连接,一类是外部网向内的网络连接,另一类是局域网之间的连接。

在网络连接测度多元组中,除了源地址和目标地址外,其余的都被认为是异常一致性因素。

5.2 建立检测模型

¹ 数据来源于 <http://iris.cs.uml.edu:8080/network.html>

本实验建立的评语集包含四种评语，对应于{“正常”，“有点异常”，“异常”，“很异常”}。将“异常”评语以上的评估划分为入侵行为。

在异常一致性因素中，除了“服务类型”和“协议执行状态”是名义尺度的测度外，其余的都是间隔尺度，即用连续的量来表示的测度。对于间隔尺度的测度，直接使用它们的概率密度函数作为检测模型是不适合的，而测度值在一定范围的出现概率才是有意义的，因此需要首先将该测度的取值范围划分为区间，然后计算区间的密度分布函数。

5.3 实验结果

本实验分为两部分，第一部分的实验不考虑算法优化，即：

- (1) 不考虑非异常一致性因素对 FJADA 的影响，即设它们对模糊综合评判结果没有影响。即与所有“异常度”评语的模糊关系隶属度为 0。
- (2) 不考虑复合因素。即忽略评价因素之间的互相影响，设它们是相互独立的因素。

第二部分的实验考虑上述两个因素对算法的影响。

5.3.1 第一部分的实验结果及其分析

	评判为入侵事件的网络连接比率		
网络连接	对外的连接	向内的连接	网间连接
Normal	0.58%	0.21%	0.85%
Intrusion1	3.05%	15.39%	18.78%
Intrusion2	3.86%	12.51%	12.53%
Intrusion3	2.58%	20.19%	9.33%

表 1 正常和异常数据的入侵行为比率

将第一部分的实验结果与 Lee 的相比较，发现检测效果总的来说要比 Lee 的算法的效果要差，主要体现在对最后三组数据的检测率不如 Lee 的算法。但从对 Normal 的数据的检测来看，其误检率(false positive rate)要低于 Lee 的相同参数，这从实验中证明了基于测度值概率密度函数的检测模型的稳定性要高于基于数据挖掘的检测模型。

5.3.2 第二部分的事件结果及其分析

	评判为入侵事件的网络连接比率		
网络连接	对外的连接	向内的连接	网间连接
Normal	0.41%	0.35%	0.25%
Intrusion1	1.74%	31.46%	40.66%
Intrusion2	2.13%	33.82%	11.08%
Intrusion3	4.15%	44.13%	15.10%

表 2 算法优化后检测到的入侵行为比率

这次实验的结果较上次有了较大的改进，最明显的变化是对后三组入侵数据的检测率有了很大的提高，从而显示了这两个因素对提高检测效果的重要性。用这次的数据同 Lee 的做比较，发现 FJADA 无论从提高检测率还是减少误检率都明显优于 Lee 的方法。

六. 总结

FJADA 算法的优点主要有两个：第一，检测模型的建立对训练数据的要求不高，只需要足够、真实的网络数据；而且检测模型一旦建立就比较稳定，不需要频繁地重新学习和建立就能达到较为满意的检测率和较低的误报率，而这也是以往的异常检测方法所缺乏的；第二，采用模糊综合评判作为检测算法，该方法的特点就是能将比较复杂、模糊性强的问题用精确的数学工具来解决，从而获得比较精确的结果，从实验结果来看，它获得了比 Lee 的方法更低的误报率和更高的检测率，这证明了该方法是有效的。

本算法还可以进一步优化，一方面可以通过调节测度权重向量和采用新的数学模型来定义测度与评语之间的模糊关系使检测性能提高，上述实验采用线形函数定义模糊关系，实际可以采用一些更为复杂的数学模型；另一方面是为网络连接定义更完善的测度集，例如一些

统计性测度，这样不但能拓宽检测面，还能提高检测性能。

FJADA 能够检测到未知的攻击方式，这里的关键在于当按照 3.4 节的攻击分类方法，属于同类的攻击方式的特征权重向量是相同的，因此就算某种攻击方式是未知的，只要它属于已定义的攻击类别也能被检测出来。

FJADA 对检测拒绝服务和扫描攻击还比较困难，因为这类的攻击不能用网络连接来描述。解决这个问题方法之一是将“在一段时间内源地址与宿地址之间的报文集合”作为另一个检测对象。

【参考文献】

1. T Lunt, A Tamaru, F Gilham, R Jagannathan, P Neumann, H Javitz, A Valdes, and T Garvey. A real-time intrusion detection expert system(IDES)-final technical report. Computer Science Laboratory, SRI International, Tech Rep: SRI-CSL-95-07, 1992.
2. W Lee, and S J Stolfo. Data mining approaches for intrusion detection. In Proceedings of the 7th USENIX Security Symposium, 1998, 21(3): 181~199.
3. T LANE, and C E Brodley. An application of machine learning to anomaly detection. In Proceedings of the 20th National Conference on National Information Systems Security. Vol1, Baltimore, MD, 1997, 366~380.
4. J E Dickerson and J A Dickerson. Fuzzy network profiling for intrusion detection. In International Conference of the North American Fuzzy Information Processing Society, Vol1, Atlanta, Georgia, USA, 2000, 301~306
5. J E Dickerson, J Juslin, O Koukousoula, and J A Dickerson. Fuzzy intrusion detection. In IFSA World Congress and 20th NAFIPS International Conference, Vol3, Vancouver, British Columbia, Canada, 2001, 1506~1510
6. A Siraj, S M Bridges, R B Vaughn. Fuzzy cognitive maps for decision support in an intelligent intrusion detection system. Vol4, In IFSA World Congress and 20th NAFIPS International Conference, Vancouver, British Columbia, Canada, 2001, 2165~2170
7. 李之棠, 杨红云. 模糊入侵检测模型. 计算机工程与科学, 2000, 22(2): 49~53
(Li Zhitang, Yang Hongyun. A fuzzy intrusion detection model. Journal of Computer Engineering and Science (in Chinese), 2000, 22(2): 49~53)
8. 何新贵. 模糊知识处理的理论和技术. 第二版. 北京: 国防工业出版社, 1998
(He Xingui. Fuzzy theories and fuzzy techniques in knowledge processing (in Chinese) .2nd Edition. Beijing: National Defense Industry Press, 1998)