

基于不满意度的网络安全模型^{*}

王海晟^{1†}, 桂小林²

(1. 西安理工大学 计算机学院, 西安 710048; 2. 西安交通大学 电子信息工程学院, 西安 710049)

摘要: 提出了一种基于不满意度的网络安全模型, 主要功能是帮助用户在网络环境中正确地选择交易对象, 屏蔽恶意节点, 基于不满意度 (degree of dissatisfaction, DoD) 对交易节点进行分类控制。节点的不满意度定义为该节点属于恶意节点集的概率。a) 使用粗糙集 (rough set) 模块与 Bayesian 学习器计算节点的不满意度, 依据节点的交易历史记录计算节点的本地不满意度 (local DoD, LDoD), 依据反馈推荐意见计算推荐不满意度 (recommended DoD, RDoD), 基于不满意度将节点划分为可信任节点、陌生节点、恶意节点等不同的类型; b) 基于推荐意见的信息熵 (information entropy) 计算其可信度, 对反馈推荐意见进行综合。实验表明, 与已有的安全模型相比, 提出的安全管理模型对恶意节点具有更高的检测率, 具有更满意的交易成功率。

关键词: 网络安全模型; 不满意度; 粗糙集; 贝叶斯学习器; 信息熵; 仿真

中图分类号: TP393.08

文献标志码: A

文章编号: 1001-3695(2013)02-0566-04

doi:10.3969/j.issn.1001-3695.2013.02.069

New network security model based on degree of dissatisfaction

WANG Hai-sheng^{1†}, GUI Xiao-lin²

(1. School of Computer Science & Technology, Xi'an University of Technology, Xi'an 710048, China; 2. School of Electronic & Information Engineering, Xi'an Jiaotong University, Xi'an 710049, China)

Abstract: This paper proposed a new security model based on the degree of dissatisfaction (DoD). The main function was to help users in network environment to select correctly the transaction object and to avoid malicious node. This paper defined the degree of dissatisfaction as the probability that a node belonged to the malicious node set. This model used the rough set and Bayesian learner to calculate the DoD of the nodes, and based on historical transaction records to calculate the local DoD (LDoD), based on feedbacks on recommendations to calculate the recommended DoD (RDoD), and divided nodes into trust nodes, strange nodes and malicious nodes. It calculated the credibility of feedbacks based on the information entropy, and used the credibility of feedbacks to integrate feedbacks. Compared with existing trust models, the experiment indicates that the model can obtain higher examination rate over malicious nodes, with the higher transaction success rate.

Key words: network security model; degree of dissatisfaction; rough set; Bayesian learner; information entropy; simulation

0 引言

为了建立适应网络动态复杂性及恶意节点行为多样性的分布式信任模型, 使信任机制更加客观合理, 已经提出了许多安全模型^[1,2]。文献[3]中提出了基于群组的网络信任管理模型, 在计算目标节点的总体信任度时, 综合了节点的直接信任度、组内信任度及组间信任度。文献[4~6]给出了基于信誉的评价模型。现有的网络安全模型都是基于信誉度或信任度等概念, 而没有对节点之间交易失败的事件进行分类量化。本文提出了一种基于不满意度对交易节点进行分类控制的实用网络安全模型。模型的重点是根据节点间的交易历史记录, 结合粗糙集模块和 Bayesian 学习器计算节点的不满意度, 将恶意节点检测出来并对其进行控制。

本文提出的安全模型既可以用于无线传感器网络 (WSN), 也可以用于 P2P 网络。本文以在 P2P 网络中的应用来描述提出的安全模型, 以对等文件共享应用为例来描述所提

出的算法。

1 不满意度的定义

本文对发生在交易节点之间的失败的交易划分为两大类: 严重失败与一般不满意。严重失败事件包括下载了恶意攻击文件、欺诈交易、大规模交易且质量严重不满意和恶意反馈推荐等; 一般不满意事件包括小规模交易且质量不满意、下载速度太低或中间离线、反馈推荐意见偏离等。

节点划分为可信任节点、陌生节点 (不确定节点) 和恶意节点三个类型。可信任节点集合用 C_{trust} 表示; 恶意节点集合用 C_{virus} 表示; 陌生节点集合用 C_{stranger} 表示。

定义1 节点 X_{ei} 属于恶意节点的概率 $P(C_{\text{virus}} | X_{ei})$ 定义为节点 X_{ei} 的不满意度 (DoD)。

定义2 依据节点 X_{ei} 与本地节点交易的历史记录, 计算得到的该节点属于恶意节点的概率定义为本地不满意度 (LDoD)。

收稿日期: 2012-05-24; 修回日期: 2012-07-10 基金项目: 国家自然科学基金资助项目 (60873071)

作者简介: 王海晟 (1978-), 男 (通信作者), 博士研究生, 主要研究方向为分布式网络计算、网络安全、并行计算 (haisheng.wang@yahoo.cn); 桂小林 (1966-), 男, 教授, 博导, 主要研究方向为网络计算、动态信任管理理论。

2 使用机器学习技术计算本地不满意度

2.1 Bayesian 学习器的训练与应用

假设 W 是给定世界的有限或无限的所有观测对象的集合, $F: X \rightarrow \{0, 1\}$, $X \in W$, 是该世界的模型。由于观察能力的限制, 只能获得这个世界的一个有限子集 $\langle x_1, d_1 \rangle \cdots \langle x_m, d_m \rangle$, 称为样本集, 其中 x_i 为 X 中的某个实例, d_i 为 x_i 的目标函数值, $d_i \in \{0, 1\}$ 。机器学习就是根据这个样本集, 推算这个世界的模型 $f: X \rightarrow \{0, 1\}$, 使得 f 是 F 的一个近似, 即 $d_i = f(x_i)$ 。

这里学习器又可称为分类器, 当强调通过训练样本集进行训练、机器学习时, 称为学习器; 当强调该学习器用于分类目的时, 称为分类器^[7-9]。Bayesian 分类器就是要在候选假设分类集合 H 中寻找当给定数据 D 时可能性最大的分类假设 $h \in H$ 。这样的具有最大可能性的假设被称为极大后验 (maximum a posteriori, MAP) 假设。确定 MAP 假设的方法是用 Bayesian 公式计算每个候选假设的后验概率。当下式成立时, 称 h_{MAP} 为 MAP 假设。

$$h_{\text{MAP}} = \arg \max_{h \in H} P(h|D) = \arg \max_{h \in H} \frac{P(D|h)P(h)}{P(D)} \quad (1)$$

使用训练样本集对 Bayesian 学习器进行训练, 这里使用的训练样本集就是本地节点与其他节点的交易历史记录。设有 m 个节点分类类型 $C_1, C_2, \dots, C_m$ 。给定一个未知的数据样本 X_{ei} (即等待分类的样本), Bayesian 学习器将首先计算该样本 X_{ei} 属于各个分类类型的概率。Bayesian 分类器将未知的样本分配给类 C_i , 当且仅当

$$P(C_i|X_{ei}) > P(C_j|X_{ei}) \quad 1 \leq j \leq m, j \neq i \quad (2)$$

其中: C_i 是具有最大后验概率的类。根据式(1), 后验概率为

$$P(C_i|X_{ei}) = \frac{P(X_{ei}|C_i)P(C_i)}{P(X)} \quad (3)$$

利用 Bayesian 学习器计算节点属于各种分类的概率, 样本数据是由节点间的交易历史记录通过汇总得到的, 每条交易记录包括若干交易特征。本文选取的特征列划分为三个部分: 其一是所有时间段内 (本文取最近 30 个时间段) 的统计数据; 其二是最近三个时间段内的统计数据; 其三是最近两个时间段内的统计数据。

截取交易历史记录汇总表中的部分内容说明如下: 设某个节点 i 的特征向量为 $X_i = (x_1, x_2, \dots, x_n)$, 其中 $x_1, x_2, \dots, x_n$ 分别是节点 i 对应 $X_1, X_2, \dots, X_n$ 特征项的取值。 $X_1, X_2, \dots, X_n$ 特征项对应于节点交易记录汇总表中的各列如表 1 所示。

表 1 节点交易记录汇总表 (训练数据集截取)

ID	所有时间段内 (最近 30 个时间段)					最近两个时间段内					分类标号
	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	C_i	
100101	...	...	...	...	...	...	...	...	...	...	C_{trust}

其中: X_1 为所有的时间段内与本地节点交易的总次数, 在运行开始阶段, 划分为 $X_1 \geq 100$ 、 $100 > X_1 \geq 50$ 、 $50 > X_1 \geq 20$ 和 $20 > X_1$ 等区间; X_2 为所有时间段内恶意攻击的总次数, 划分为 $X_2 \geq 4$ 、 $X_2 = 3$ 、 $3 > X_2 \geq 1$ 和 $X_2 = 0$ 等区间; X_3 为所有时间段内欺诈交易的总次数, 划分为 $X_3 \geq 5$ 、 $5 > X_3 \geq 3$ 、 $3 > X_3 \geq 1$ 和 $X_3 = 0$ 等区间; X_4 为所有时间段内恶意反馈总次数, 划分为 $X_4 \geq 6$ 、 $6 > X_4 \geq 4$ 、 $4 > X_4 \geq 1$ 和 $X_4 = 0$ 等区间; X_5 为最近两个时间段

内恶意攻击的次数; 划分为 $X_5 \geq 3$ 、 $X_5 = 2$ 、 $X_5 = 1$ 、 $X_5 = 0$ 等区间; X_6 为最近两个时间段内欺诈交易的次数, 划分为 $X_6 \geq 4$ 、 $X_6 = 3$ 、 $3 > X_6 \geq 1$ 、 $X_6 = 0$ 等区间; X_7 为最近两个时间段内恶意反馈总次数, 划分为 $X_7 \geq 5$ 、 $5 > X_7 \geq 3$ 、 $3 > X_7 \geq 1$ 、 $X_7 = 0$ 等区间; X_8 为最近两个时间段内其他交易失败的次数; X_9 为最近时间段内交易的失败率, $X_9 = nf/n$, n 是最近时间段内本地节点与该节点交易的总次数, nf 是交易失败的总次数。

实验初始训练样本包括本地节点与 240 个交易节点的交易记录。其中可信节点 60 个, 恶意节点 60 个, 陌生节点 120 个。通过对训练集学习而归纳出学习器, 本系统中称这个学习器为 Bayesian 学习器 A。

对未知样本 X_{ei} , 依据该节点与本地节点交易的历史记录, 使用 Bayesian 学习器分别计算该节点属于可信节点的概率 $P(C_{\text{trust}}|X_{ei})$ 、属于陌生节点的概率 $P(C_{\text{stranger}}|X_{ei})$ 和属于恶意节点的概率 $P(C_{\text{virus}})$ 。

2.2 粗糙集模块的训练与应用

本文用四元组 $S = (U, A, V, f)$ 作为一个知识表达系统。其中, U 为对象的有限集合, 称为论域; A 为非空有限集称为属性集合; V 为属性 A 的值域; $f: U \times A \rightarrow V$ 是一个信息函数, U 中任一元素其属性 A 取值 a 时, a 在 V 中有唯一确定值。令 R 为一族等价关系, $r \in R$, 如果 $\text{ind}(R) = \text{ind}(R - \{r\})$, 即属性集 R 对于对象集 U 的分类与属性集 $R - \{r\}$ 对于 U 的分类相同, 则称 r 为 R 中冗余; 否则 r 为 R 中必要。若存在 $Q = P - r$, $Q \subseteq P$, Q 是独立的, 满足 $\text{ind}(Q) = \text{ind}(P)$, 则称 Q 为 P 的一个约简, 用 $\text{red}(P)$ 表示。一族等价关系 P 可能有多个约简, 全部约简的交集定义为 P 的核, 记为 $\text{core}(P)$, $\text{core}(P) = \bigcap \text{red}(P)$ 。约简是在不丢失信息的前提下, 能够与原信息系统表达同样知识的最小条件属性集, 它是保持信息系统的相同分类能力的最简形式, 通过知识约简导出问题的分类规则。

Bayesian 学习器计算一个节点属于某个分类类型的概率时, 所有的属性都起作用。计算节点的不满意度时, 即计算节点属于恶意节点的概率时, 经常有仅根据一个或两个属性就可以确定节点属于恶意节点概率的情况。使用粗糙集模块可以快速、准确地、仅根据一个或两个属性确定节点属于恶意节点的概率。

使用对 Bayesian 学习器进行训练时使用过的、相同的训练样例集, 无须提供训练集以外的任何先验信息, 粗糙集就可以生成精确的、可检查的分类规则。结合使用粗糙集与 Bayesian 学习器来完成对节点的不满意度的计算, 提高了计算准确度和效率。处理过程描述如下:

a) 使用在上一节描述的对 Bayesian 学习器 A 进行训练时使用的 240 个训练样本集 (含有 240 个记录的交易记录表就是一个决策表), 对粗糙集模块进行训练, 通过离散化、属性约简、属性值约简导出精确而又易于检查和证实的分类规则。

b) 按照如下要求, 选择、保留部分分类规则, 生成粗糙集模块 A:

(a) 只选择、保留可信度为 100% 的分类规则。

(b) 只选择、保留决策属性值对应于恶意节点的分类规则。

(c) 只选择、保留仅包含一个或两个条件属性的分类规

则,比如,if(所有时间段内恶意攻击次数(X_2) ≥ 5) then(该节点分类为恶意节点);或者 if(所有时间段内恶意攻击次数(X_2) $=3$ and 所有时间段内恶意反馈次数(X_4) $=3$) then(该节点分类为恶意节点)等。

c)仅当训练样本集改变时,重新生成与选择保留分类规则。

d)对交易节点进行分类判断,即对未知样本 X_{ei} 进行分类判断时,进行以下步骤:

(a)调用粗糙集模块,使用选择保留的分类规则对待分类节点进行分类,如果分类成功,则设置该节点属于恶意节点的概率 $P(C_{virus}|X_{ei}) = 100\%$ 。

(b)如果使用粗糙集模块分类不成功,不能识别,那么调用 Bayesian 学习器,分别计算该节点属于可信节点的概率 $P(C_{trust}|X_{ei})$ 、属于陌生节点的概率 $P(C_{stranger}|X_{ei})$ 和属于恶意节点的概率 $P(C_{virus}|X_{ei})$ 。属于恶意节点的概率 $P(C_{virus}|X_{ei})$ 就是节点的本地不满意度 $LDoD = P(C_{virus}|X_{ei})$ 。

e)每当交易记录的时间段交替时或交易中出现了严重失败事件后,立即调用粗糙集与 Bayesian 学习器对相关节点进行检测,计算该节点属于恶意节点的概率等。

3 推荐不满意度

3.1 使用粗糙集和 Bayesian 学习器计算 RDoD

反馈推荐意见是按照表 1 的格式给出相关的一行记录,包括交易节点标志号 ID(这里的交易节点标志号 ID 就是被推荐节点的 ID)、所有时间段内与被推荐节点的交易总次数 X_1 、所有时间段内发生的恶意攻击总次数 X_2 、欺诈交易总次数 X_3 、恶意反馈推荐次数 X_4 、最近时间段内恶意攻击次数 X_5 、欺诈交易次数 X_6 、恶意反馈推荐次数 X_7 、其他交易失败次数 X_8 、最近时间段内交易失败率 X_9 和推荐的该被推荐节点所属的节点类型标号。通常可以收到许多节点的推荐意见,组成一个推荐意见综合表。通过对由反馈推荐意见表综合而成的训练样例集学习而归纳出学习器,本系统中称这两个学习器为粗糙集模块 B 和 Bayesian 学习器 B。

对被推荐节点 X_{ei} ,依据推荐意见综合表,使用粗糙集模块 B 和 Bayesian 学习器 B 分别计算被推荐节点属于可信节点的概率 $P(C_{trust}|X_{ei})$ 、属于陌生节点的概率 $P(C_{stranger}|X_{ei})$ 和属于恶意节点的概率 $P(C_{virus}|X_{ei})$ [10-12]。

定义 3 推荐不满意度(RDoD)。依据反馈推荐意见计算得到的被推荐节点 X_{ei} 属于恶意节点的概率。

3.2 基于信息熵的可信度计算

熵的概念来源于热力学。在热力学中熵的定义是系统可能状态数的对数值,称为热熵。信息熵是由香农(Shannon)将热力学熵引入信息论而提出的。熵是体系的混乱度或无序度的度量。在信息论中信源输出是随机量,因而其不定度可以用概率分布来度量。信息熵 H 的计算公式为

$$H(X) = H(p_1, p_2, \dots, p_n) = -\sum p(x_i) \log p(x_i)$$

其中: $p(x_i)$ ($i=1, 2, \dots, n$) 为信源取第 i 个符号的概率。

在对推荐意见进行综合的过程中,推荐意见一致时,可信度高,推荐意见分散时,可信度低。推荐意见一致时,信息熵就低;推荐意见分散时,推荐意见越是混乱,信息熵就越高。可

见,对于本文讨论的情况,推荐意见的可信度与推荐意见的信息熵成反比。

计算推荐意见的信息熵实质是计算推荐意见的分散度,计算被推荐节点可能属于哪个节点类型的不确定度。推荐意见分为三种情况:被推荐节点可能被推荐为可信节点、陌生节点或恶意节点。依据推荐意见表分别计算推荐被评价节点为可信节点、陌生节点与恶意节点的推荐节点数;分别计算该被推荐节点属于可信节点、陌生节点与恶意节点的概率,如表 2 所示。

表 2 对一个被推荐节点的推荐意见综合分析示例

比较项	推荐被评价节点 为可信节点	推荐被评价节点 为陌生节点	推荐被评价节点 为恶意节点
推荐节点数	70	20	10
百分比	70%	20%	10%
信息熵 H	$H = -(0.7 \log_2 0.7 + 0.2 \log_2 0.2 + 0.1 \log_2 0.1) = 1.157$		

依据所有节点的反馈推荐构成的推荐表 T_{all} ,得到所有节点的推荐意见的信息熵为 H_{all} 。所有节点的推荐意见的可信度定义为 $T_{all} = 1/H_{all}$, T_{all} 最大取 1.0。

4 基于综合不满意度 DoD 对节点进行分类管理

节点的综合不满意度 DoD 按照以下规则进行计算:

a)节点的综合不满意度 DoD 取节点的本地不满意度的值:

$$DoD = LDoD \quad \text{当 } LDoD > 0.6$$

b)不满足上述条件时,节点的综合不满意度 DoD 取节点的推荐不满意度的值:

$$DoD = RDoD \quad \text{当 } RDoD > 0.6$$

c)不满足上述两组条件时,节点的综合不满意度 DoD 按照下式计算:

$$DoD = LDoD / (1 + T_{all}) + T_{all} \times RDoD / (1 + T_{all})$$

其中: T_{all} 是对被评价节点的全体推荐意见的可信度,在上一节中根据推荐意见的信息熵计算获得。

定义 4 节点 X_{ei} 的不满意度 DoD 大于 0.6 时,被认定为恶意节点。

建立恶意节点列表,在若干时间段内不与恶意节点发生交易,拒绝接收这些节点的反馈推荐意见。在仿真实验中,一旦发生严重失败事件,立即更新相应的交易记录汇总表,并调用粗糙集模块和 Bayesian 学习器,计算该节点的本地不满意度,判断由于这一严重失败事件该节点是否应该重新分类。如果检测出恶意节点,立即添加到恶意节点列表。需要与陌生节点进行交易时,先请求其他节点对该陌生节点进行推荐,调用粗糙集模块和 Bayesian 学习器,计算该节点的推荐不满意度,确定是否可以与该陌生节点进行交易。

5 仿真及其结果分析

本文通过在 PeerSim 平台上集成了粗糙集模块和 Bayesian 学习器,实现了一个模拟对等网络环境来对本文的相关模型及其算法进行性能分析 [13,14]。在仿真实验中,对等网络的规模为 1 000 个节点,在每个仿真周期中每个节点至少与 30 个以上的节点进行交互。在实验中构建了两种安全模型:

a)基于不满意度的网络安全模型,系统构建了粗糙集分类器和 Bayesian 分类器,根据不满意度对节点进行分类管理

(本文模型);b)没有使用分类器,没有对节点进行分类控制,基于交易成功率计算节点的信誉度,按照节点的信誉度从高到低选择文件提供者(比较模型)。

实验的目的是检验基于不满意度的网络安全模型是否能够更加有效地帮助用户找到合适的服务提供者,保证提供服务的成功率。用文件服务的成功率来反映安全模型的能力。设在节点交互过程中某个时刻 t 检测到提供文件服务的总次数为 $S(t)$,提供文件服务成功的次数为 $N(t)$,则文件服务的成功率($SR(t)$)定义为 $SR(t) = N(t)/S(t)$ 。系统初始设定的恶意节点所占的百分比为 β ,该值直接影响仿真实验初始的成功率。如果 $\beta = 20\%$,仿真实验初始的成功率应该为 $1 - 20\% = 80\%$ 左右;随着系统对恶意节点的检测与控制,文件服务的成功率从80%不断提高。

通过实验检查系统的动态适应性。动态适应性就是系统在各种不确定因素的影响下提供可靠服务的能力。在仿真实验中,考虑了两个方面的动态变化因素:a)部分节点可能随时离开或加入;b)部分恶意节点可能频繁地变换行为方式,时而提供恶意服务,时而提供正常服务。

a)提供恶意文件下载服务的节点比率 mf ;b)提供恶意反馈评价意见的节点比率 mb ;c)可能随时离开或加入的节点比率 mj ;d)动态变换恶意行为方式的节点比率 md 。

实验中采用了三种设置。设置 a) $mf = 0.3$; $mb = 0.3$; $mj = 0.3$; $md = 0.3$; b) $mf = 0.25$; $mb = 0.2$; $mj = 0.2$; $md = 0.3$; c) $mf = 0.2$; $mb = 0.3$; $mj = 0.1$; $md = 0.2$ 。

图1、2是恶意节点的比例分别为总节点数的20%与30%时的文件服务成功率统计。从图中可看到,开始时两种算法的性能差距不大,但随着迭代次数的增加,基于不满意度的网络安全模型显示出具有更高的成功率。

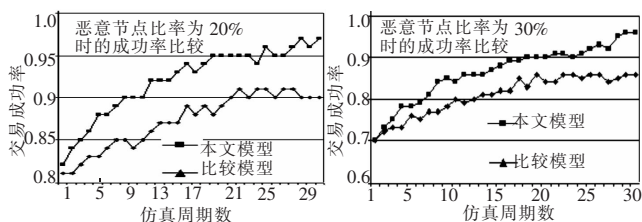


图1 交易成功率 $SR(t)$ 比较1 ($\beta=20\%$)

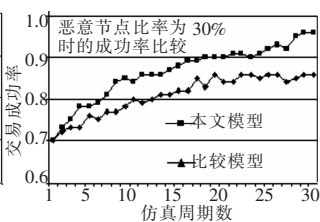


图2 交易成功率 $SR(t)$ 比较2 ($\beta=30\%$)

依据图3仿真结果可以得出,本文的模型与基于高信誉度优先选择的模型相比较,本文的模型在网络动态变化的情况下和较高的文件服务交互成功率,具有更好的网络动态适应性。这是因为本文的模型对交互失败的事件进行了更详细的分析和更细化的分类统计,并对恶意节点进行了有效控制。

图4是本文算法在各种不确定因素变化的情况下动态性能的纵向比较,可以看出在恶意节点比例增加的情况下(恶意节点比率分别为20%、25%和30%),本文算法的性能依然可以保持较高的水平,交易成功率都逐渐达到90%以上,具有令人满意的动态适应性。

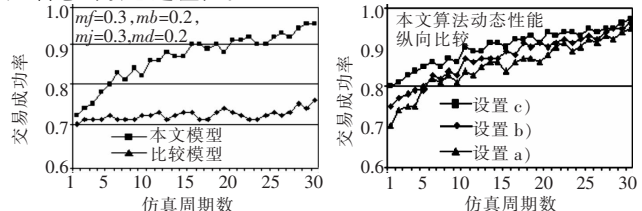


图3 动态适应性比较

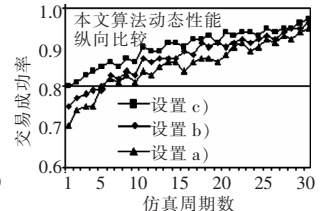


图4 本文算法动态性能检测

6 结束语

本文提出了一种基于不满意度的网络安全模型。该模型放弃了局部信任度与全局信任度等概念,采用对不满意事件进行分类统计,对交易节点进行分类控制。模型中给出了不满意度的定义、计算和应用,可以帮助用户正确选择交易对象,屏蔽恶意节点。实验表明,与已有的安全模型相比,当网络中恶意节点的比例为总节点数20%、30%时,本文模型可以提高交易成功率10%~15%。本文模型是可行、有效的。

参考文献:

- [1] SUN Y L, HAN Zhu, YU Wei, et al. A trust evaluation framework in distributed networks: vulnerability analysis and defense against attacks [C]//Proc of the 25th IEEE Conference on Computer Communications. 2006:230-236.
- [2] 姜守旭, 李建中. 一种P2P电子商务系统中基于声誉的信任机制[J]. 软件学报, 2007, 18(10): 2551-2563.
- [3] 施荣华, 辛晶晶. 一种基于群组的P2P网络信任管理模型[J]. 计算机应用研究, 2010, 27(7): 2638-2640.
- [4] SELCUK A A, UZUN E, PARIENTE M R. A reputation-based trust management system for P2P networks[J]. International Journal of Network Security, 2008, 6(3): 235-245.
- [5] MORDACCHINI M, BARAGLIA R, DAZZI P, et al. A P2P recommender system based on gossip overlays PREGO [C]//Proc of the 10th IEEE International Conference on Computers and Information Technology. 2010:83-91.
- [6] SRIVATSA M, LIU L. Securing decentralized reputation management using TrustGuard [J]. Journal of Parallel and Distributed Computing, 2006, 66(9): 1217-1232.
- [7] TIAN Jun-feng, TIAN Rui. A fine-grain trust model based on domain and Bayesian network for P2P e-commerce system [J]. Journal of Computer Research and Development, 2011, 48(6): 974-982.
- [8] CHEN Yu, PEDRYCZ W, WATADA J. A fuzzy regression based support vector machine(SVM) approach to fuzzy classification[J]. ICIC Express Letters, 2010, 4(6): 2355-2362.
- [9] LIM S, KEUNG C, GRIFFITHS N. Towards improved partner selection using recommendations and trust [C]//FALCONE R, et al. Trust in Agent Societies (TRUST), Lecture Notes in Computer Science. Berlin: Springer, 2008:43-64.
- [10] 张骞, 张霞, 文学志, 等. Peer-to-peer环境下多粒度Trust模型构造[J]. 软件学报, 2006, 17(1): 96-107.
- [11] YU Zhi-hua. Analysis of malicious behaviors in peer-to-peer trust model [J]. Computer Engineering and Applications, 2007, 43(13): 18-21.
- [12] GUO Xiao-qiong, GUAN Hai-bing. A trust management model based on Bayesian network [J]. Information Security and Communications Privacy, 2008(2).
- [13] 赵越, 茹婷婷. 贝叶斯网络结构学习方法新探[J]. 长春大学学报, 2011, 21(6): 32-34.
- [14] CHENG C F, WANG S C, LIANG T. File consistency problem of file-sharing in peer-to-peer environment [J]. International Journal of Innovative Computing Information and Control, 2010, 6(2): 601-613.