

基于声明机制的 Web 信任模型^①

邓忠军^{②*} 王少杰^{**} 郑雪峰^{*} 林 冉^{***} 锁延锋^{*} 于 真^{*}

(^{*} 北京科技大学信息工程学院 北京 100083)

(^{**} 国家信息技术安全研究中心 北京 100084)

摘 要 为了解决 Web 服务中的节点评价存在节点的计算性能和网络连接等方面的物理差异、节点评价标准的主观兴趣差异及信任评价考虑因素的全面性差异的问题,提出了一种基于声明机制的信任评价模型。该模型通过建立声明机制的方法,考虑了节点物理差异和主观差异,对信任评价标准进行了细化。仿真实验分析表明,新模型精化了信任算法的粒度,提高了信任评价的准确度,体现了节点的个性化特性,具有很好的可扩展性。

关键词 Web 服务, 声明机制, 信任

0 引 言

利用 Web 信任模型度量网络中节点间的信任关系是目前引起广泛关注的网络安全研究的热点之一。2001 年 Jøsang 提出了用事实空间和概念空间来描述和度量信任关系^[1],提供了用于信任度的推导和综合计算的逻辑算子,但对信任关系的动态变化尤其是信任关系的反馈更新等方面存在不足,无法有效地抵御恶意推荐信息。2003 年徐锋等提出了一个用于度量软件服务间信任关系的信任评估模型^[2],该模型将信任抽象成一个由信任评估主体对客体的主观期望和客观经验共同作用的函数,还提供了一个用于综合直接经验和第三方推荐经验的合理方法。2004 年 Xiong 等给出了一个适用于 P2P 电子社区的局部信任模型^[3],该模型将节点的可信度视为对以往该节点向其它节点提供服务的水平的综合评价。此模型引入了节点对交互的反馈、反馈的可信度、节点参与交互的次数、交互的属性和节点所在社区等 5 个因素来度量节点的可信程度。此外,还有其它一些模型,包括 EB-DTM^[4]、多粒度 Trust 模型^[5]、DyTrust^[6]等。这些模型提供了描述、量化、传递及综合信任信息的功能,且对信任信息的操作均以节点间的推荐信任关系为基础,能够在一定程度上解决信任评价的准确性问题。考虑到分布式网络的复杂性,我们认为在进行信任评价时还需要考虑分布式网络中提供服务的节点存在物理差异、访问

节点的评价标准存在差异这两个主要因素。现有的信任模型可以通过基于行为相似度的算法,分配节点信任评价值的权重,但当对大规模网络的访问非常稀疏时,基于行为相似度算法难以评价两个节点是否真的具有行为上的相似性,而陌生人之间通过“自我介绍”进行沟通以促成合作的方式在某种程度上更容易解决没有直接交互历史的问题,利于提高合作的效率和降低风险。同时也要看到,经常更改“自我介绍”的人会让人感到不可信。借鉴“自我介绍”这一行为,本文提出了一种新的基于声明机制的 Web 服务信任模型。该模型考虑了节点的物理差异和评价标准差异,对信任评价标准进行了细化。仿真实验分析表明,此模型精化了信任算法的粒度,提高了信任评价的准确度,体现了节点的个性化特性,具有很好的可扩展性。

1 声明机制

本文所采用的声明指的是:每个 Web 服务在进行服务发布时,必须对自身服务的安全性、服务的质量和服務的可靠性等方面进行声明,表明自己的服务性能和参与网络的立场和态度。

首先假设 Web 服务的信任评价主要考虑服务的安全性、服务的质量和服務的可靠性这三个方面,同时为了简化问题,仅考虑这三方面性能中的病毒及恶意插件、平均数据包下载速度、平均服务的响应

^① 863 计划(2007AA012474),国家发改委信息安全专项([2008]1736 号)资助项目。

^② 男,1963 年,博士生;研究方向:网络安全,信任算法,数据挖掘等;联系人, E-mail: deng.zj@163.com
(收稿日期:2009-07-10)

时间、平均无故障响应时间、提供服务的时间段这 5 个因素,如图 1 所示。

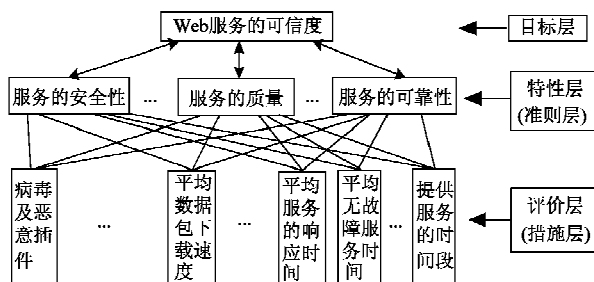


图 1 Web 服务的可信度评价

定义 1 六元组 $S^x = (X^0, X^A, X^B, X^C, X^D, X^E)$ 是节点 X 的声明集,其中: X^0 表示节点 X 能提供的服务类型; X^A 表示病毒及恶意插件,其取值为 0 或 1,取值为 0 时代表有病毒或恶意插件,即服务是不安全的,取值为 1 时代表无病毒和恶意插件,即服务是安全的; X^B 表示数据包的下载速度,其取值为正实数,是指服务所提供资源的下载速度(单位是 bps), X^B 的取值越大,代表服务的质量越高; X^C 表示服务响应时间,其取值为正实数,是访问服务时需要等待的时间(单位是 ms), X^C 的数值越小,服务质量就越高; X^D 表示无故障响应时间,其取值为一段时间间隔,指的是服务相邻两次故障之间的平均工作时间, X^D 的数值越大,服务的可靠性越高; X^E 表示提供服务的时间段,其取值为某个时间段,指的是节点在某个时间段提供服务,而在其它时间段不提供服务。

节点对服务的声明发布出来后,任何节点可通过 Web 服务注册中心查看声明。声明发布后,节点不能擅自更改声明中的内容,若要更改,需经 Web 服务注册中心的同意,同时保留原声明,其它节点在查看有过更改的服务声明时,不仅能看到该节点当前的声明,同时能看到该节点的历史声明。我们认为,经常更改服务声明的节点,在某种意义上会给服务的访问者和合作者一种不可信的感觉,这如同一个经常更改自身行为原则的人,难以让人信任。节点就自身所提供的服务进行的声明,其实也是对自身能力的一种评估,它直接影响到服务的信任度和合作成功的可能性。节点通过声明机制寻求合作或查找服务时,更容易找到个体所需要的节点或最有帮助获得成功的节点,从而避免分布式网络中与大量陌生节点进行交互时选择合作节点的盲目性,提高了服务的效率。但就现实而言,让每个节点都提供真实的信息是不太可能的,因此,在声明机制的基

础上,我们建立了基于声明机制的信任模型。

2 基于声明机制的信任模型

Web 服务中要求每个节点都发表真实的声明, Web 服务注册中心根据各个节点的声明,把同一类服务按照下载速度、服务响应时间等方面把服务分为几层。以下载速度为例,可按表 1 分为若干层次。访问节点根据自己的下载速度,选择适合自己的服务,例如通过拨号上网,网络速度只能达到 64kbps 的节点,只需访问第 5 层或第 4 层中表现良好的节点就可以了,不需要一定访问第 1 层中的节点。这种分层的方法,能够起到网络分流的作用。

表 1 声明 j^B 的 Web 服务层次表

层次	下载速度(v)	发布声明的 Web 服务节点
1	$v \geq 1\text{Mbps}$	$j_{m1}, j_{m2}, \dots$
2	$512\text{kbps} < v < 1\text{Mbps}$	$j_{n1}, j_{n2}, \dots$
3	$128\text{kbps} < v \leq 512\text{kbps}$	$j_{s1}, j_{s2}, \dots$
4	$64\text{kbps} < v \leq 128\text{kbps}$	$j_{y1}, j_{y2}, \dots$
5	$v \leq 64\text{kbps}$	$j_{z1}, j_{z2}, \dots$

2.1 基于声明的信任评价

假设节点 i 是一个访问节点, i 的声明集为六元组 $S^i = (i^0, i^A, i^B, i^C, i^D, i^E)$, 当节点 i 需要得到某类服务时,首先根据自身的需求对 Web 服务注册中心进行查询,假如节点 i 最看重的是 i^B 这个因素,那么通过 Web 服务注册中心的查询结果将得到表 1 所示的层次表,节点 i 根据自身的 i^B 值选择与第几层的服务发布者进行合作或交互。

假设节点 j 能够提供某个 Web 服务 j^0 , 同时 j 发布的声明集为 $S^j = (j^0, j^A, j^B, j^C, j^D, j^E)$, 节点 i 根据节点 j 的声明,对 j 提供的服务进行访问,直接交互后得到的结果表示为六元组 $d_{ij} = (j^0, d_{ij}^A, d_{ij}^B, d_{ij}^C, d_{ij}^D, d_{ij}^E)$, 其中 j^0 表示 j 所提供服务的种类; d_{ij}^A 表示节点 j 提供的服务中是否有病毒和恶意插件,取值为 0 表示有,1 表示没有; d_{ij}^B 表示节点 i 下载节点 j 所提供资源的实际下载速度; d_{ij}^C 表示节点 i 等待节点 j 所提供服务的实际等待时间; d_{ij}^D 表示节点 j 对节点 i 提供服务时的实际无故障时间; d_{ij}^E 表示节点 i 访问 j 的时间段内。

节点 i 根据直接交互结果 d_{ij} , 对节点 j 所提供的服务进行信任评价,评价用六元组 $D_{ij} = (j^0, D_{ij}^A, D_{ij}^B, D_{ij}^C, D_{ij}^D, D_{ij}^E)$ 表示,其中 $D_{ij}^A, D_{ij}^B, D_{ij}^C,$

D_{ij}^D, D_{ij}^E 为节点 i 根据自身对服务的要求及对信任评价的理解,对 j 所提供服务的这五个方面的直接交互满意度评价。

定义 2 D_{ij}^A 表示节点 i 对 j 所提供服务的真实性(有无病毒和恶意插件)的信任评价价值:

$$D_{ij}^A = \begin{cases} 100\%, & d_{ij}^A = j^A \\ 0, & d_{ij}^A \neq j^A \end{cases} \quad (1)$$

这里只要节点 i 得到的结果与 j 的声明相符,我们就认为节点 j 的 j^A 声明值是可信的,当然 j 可以声明自己所提供的服务有病毒和恶意代码。但理性的节点在选择服务时,应该不会选择声明有病毒和恶意代码的节点服务。

定义 3 D_{ij}^B 表示节点 i 对 j 所提供服务的下载速度的信任评价价值,表示为

$$D_{ij}^B = \begin{cases} d_{ij}^B / i^B, & i^B \leq j^B \text{ 且 } d_{ij}^B \leq i^B \\ 100\%, & i^B \leq j^B \text{ 且 } i^B \leq d_{ij}^B \\ d_{ij}^B / j^B, & j^B \leq i^B \text{ 且 } d_{ij}^B \leq j^B \\ 100\%, & j^B \leq i^B \text{ 且 } j^B \leq d_{ij}^B \end{cases} \quad (2)$$

当 $i^B \leq j^B$ 时,意味着节点 i 的物理连接速度客观上小于节点 j 的速度,此时节点 i 对 j 所提供下载速度的信任评价,应该以自身的网络速度为依据。对于 $i^B \leq j^B$ 且实际下载速度 $d_{ij}^B \geq i^B$ 的情况,这里存在着矛盾,但考虑到现实情况下,节点 i 可能存在对自己下载速度评价不准确的问题,因此这种情况也可以接受;当 $j^B \leq i^B$ 时,意味着节点 i 的物理连接速度客观上大于节点 j 的速度,节点 i 可能会对下载速度有更高的预期值,但由于我们以节点发布的声明为信任评价的参考标准,实际下载速度跟声明中的速度越相符,其下载速度的信任值就越高。当 $j^B \leq i^B$ 且实际下载速度 $d_{ij}^B \geq j^B$ 时,我们认为节点 j 存在对自己下载速度评价不准确的问题,这种情况也可以接受。

定义 4 D_{ij}^C 表示节点 i 对 j 所提供服务的响应时间的信任评价价值,表示为

$$D_{ij}^C = \begin{cases} j^C / (d_{ij}^C - i^C), & (d_{ij}^C - i^C) \geq j^C \\ 100\%, & (d_{ij}^C - i^C) < j^C \end{cases} \quad (3)$$

定义 5 D_{ij}^D 表示节点 i 对 j 所提供服务的无故障响应时间的信任评价价值,表示为

$$D_{ij}^D = \begin{cases} d_{ij}^D / j^D, & d_{ij}^D \leq j^D \\ 100\%, & d_{ij}^D > j^D \end{cases} \quad (4)$$

定义 6 D_{ij}^E 表示节点 i 对 j 所提供服务的时段的信任评价价值,表示为

$$D_{ij}^E = \begin{cases} 100\%, & d_{ij}^E \in j^E \text{ 时 } j \text{ 提供服务} \\ 0, & d_{ij}^E \notin j^E \text{ 时 } j \text{ 提供服务} \end{cases} \quad (5)$$

节点 i 得到节点 j 的信任评价价值 $D_{ij} = (j^0, D_{ij}^A, D_{ij}^B, D_{ij}^C, D_{ij}^D, D_{ij}^E)$ 后,根据自身对服务的需求和理解,同时结合自己的经验,对节点 j 的服务进行直接信任整体评价,评价价值 WD_{ij} 表示为

$$WD_{ij} = (D_{ij}^A D_{ij}^B D_{ij}^C D_{ij}^D D_{ij}^E) (w_A w_B w_C w_D w_E)^T \quad (6)$$

其中 w_A, w_B, w_C, w_D, w_E 分别表示节点 i 对 $D_{ij}^A, D_{ij}^B, D_{ij}^C, D_{ij}^D, D_{ij}^E$ 采用的权重,满足 $0 \leq w_A, w_B, w_C, w_D, w_E \leq 1$, 且 $w_A + w_B + w_C + w_D + w_E = 1$, 其中 w_A, w_B, w_C, w_D, w_E 的值可以通过如下的方法获取: (1) 用户自行指定权重值; (2) 领域专家给出权重值; (3) 借鉴 Saaty 等人^[7,8] 提出的成对比较的层次分析方法,把相关因素两两相互对比,对比时采用相对尺度,尽可能地减少性质不同的诸因素相互比较的困难,提高准确度。这样通过式(6),节点 i 可以得到提供 j^0 类服务的节点 j 的直接信任整体评价,对于提供 j^0 类服务的节点 $j_1, j_2, j_3, j_4, j_5, \dots$ 节点 i 与他们进行直接交互后,也可以得到对这些节点的直接信任整体评价,表示为 $WD_{ij_1}, WD_{ij_2}, WD_{ij_3}, WD_{ij_4}, WD_{ij_5}, \dots$ 节点 i 可以根据这些服务节点的直接信任整体评价,决定下次跟哪个节点合作,同时也可以再考虑其它节点对服务声明的准确性评价的因素,对服务节点进行综合信任评价。每次交互后,节点 i 向存放信任评价信息的档案点^[9] 提交对 j 的信任评价价值 D_{ij} , 而不是 WD_{ij} , 其原因是节点 i 在进行整体信任评价获取 WD_{ij} 值时,考虑了自身的主观因素,别的节点很难理解这样的评价。

2.2 声明的准确性验证

对节点的声明进行准确性验证,也就是对节点的声明进行信任评价。假设节点 i 准备访问提供某类服务的节点 j , 除参考自身与 j 的直接交互经验外,还可参考其它与 j 有过直接交互的节点 $k_n (n = 1, 2, \dots)$ 对 j 的直接交互后的信任评价价值 $D_{k_n j} (n = 1, 2, \dots)$, 如表 2 所示。

表 2 访问节点对 j 的信任评价表

节点 j 的声明	$j^0, j^A, j^B, j^C, j^D, j^E$	对 j 声明的 信任评价	$j^0, T^A, T^B, T^C, T^D, T^E$
节点 k_n 的声明	$k_n^0, k_n^A, k_n^B, k_n^C, k_n^D, k_n^E$	k_n 对 j 的 信任评价	$j^0, D_{k_n j}^A, D_{k_n j}^B, D_{k_n j}^C, D_{k_n j}^D, D_{k_n j}^E$

根据表 2 可以对节点 j 所发表的声明的准确性进行评价,对于表示有病毒或恶意插件的声明 j^A , j^A 的准确性 T^A 可用式

$$T^A = \sum_{x=1}^n \frac{D_{kxj}^A}{n} \quad (7)$$

进行验证。这里采用求和平均的方法判断 j^A 的准确性, T^A 的值越大,说明节点 j 所发表的声明 j^A 越符合实际。

为验证节点 j 所提供 Web 服务资源的下载速度,我们采用声明中 $k_x^B \geq j^B$ 的节点对下载速度的信任评价,这是因为网络的实际下载速度不仅跟提供服务的节点网络带宽有关,而且与访问节点的带宽有关,采用声明中 $k_x^B \geq j^B$ 的节点的评价,更能客观地评价提供 Web 服务节点 j 对 j^B 声明的准确性。设节点 $k_x (x = 1, 2, \dots, u, \dots, m)$ 满足 $k_x^B \geq j^B$, 此时对于声明 j^B 的准确性 T^B 可用式

$$T^B = \sum_{x=1}^m \frac{C_{kxj}^B \times D_{kxj}^B}{\sum_{x=1}^m C_{kxj}^B} \quad (8)$$

进行验证。其中 C_{kxj}^B 表示节点 k_x 对节点 j 下载速度的信任评价的可信度,亦即采用节点 k_x 给出的信任评价值的权重,定义为

$$C_{kxj}^B = \sum_{u=1}^m \frac{\phi_{kukx}^j C_{kukx}^j}{\sum_{u=1}^m \phi_{kukx}^j} \quad (9)$$

式(9)满足 $k_u \neq k_x$, ϕ_{kukx}^j 表示相对经验因子 $\phi_{kukx}^j = 2^{-|m_o/m_{kuj} - m_o/m_{kxj}|}$, m_{kuj} , m_{kxj} 分别表示节点 k_u , k_x 与服务节点 j 直接交互的次数; m_o 表示节点是否有经验的门限值,当节点的经验大于 m_o 时,表示该节点有经验,否则表示该节点经验不足; C_{kukx}^j 表示节点 k_u 对 k_x 推荐 j 服务的个体反馈可信度评价,表示为

$$C_{kukx}^j = \begin{cases} C_{kukx}^j + \alpha \cdot \phi_{kukx}^j \cdot 2^{-2\Delta D^B/\theta}, & d_{kukx}^j < \theta \\ C_{kukx}^j - \beta \cdot \phi_{kukx}^j \cdot 2^{-(1-\Delta D^B/2\theta)}, & \text{其它} \end{cases} \quad (10)$$

式(10)满足 $k_u \neq k_x$, C_{kukx}^j 的取值范围是 $[0, 1]$, 当 $C_{kukx}^j < 0$ 时,取值为 0; 当 $C_{kukx}^j > 1$ 时,取值为 1。 α 和 β 分别为反馈可信度增加、减少因子的值,满足 $0 < \alpha < \beta < 1$, 一个节点提供善意推荐时,反馈可信度值的增加比较缓慢,若提供恶意推荐,可信度值将快速降低。 θ 为访问节点对推荐节点能容忍的最大评价偏差; d_{kukx}^j 表示为 $d_{kukx}^j = 2^{-(m_o/m_{kuj} - m_o/m_{kxj})} \cdot \Delta D^B$, ΔD^B 由式

$$\Delta D^B = |D_{kuj}^B - D_{kxj}^B| \quad (11)$$

计算,表示节点 k_u 和 k_x 对 j 的下载速度的评价差异。

通过式(8),可以对声明 j^B 的准确性 T^B 进行验证, T^B 的值越大,说明 j 对下载速度的声明 j^B 越准确,也就越可信。同时,在对下载速度的声明进行信任性验证时,我们考虑了个体节点经验差异的问题,通过引入相对经验因子的方法,对反馈可信度进行了更新,更能体现节点的个性化特性,提高了信任评价的准确性。

同理,我们可以对节点 j 的 j^C , j^D 和 j^E 声明的准确性进行验证,验证结果分别表示为 T^C , T^D , T^E , 验证的方法可以采用文献[5,6]等信任算法,针对不同的声明的特点,采用不同的信任评价算法,以提高信任评价的准确性。限于篇幅,具体算法不予讨论。

通过如上的准确性验证算法,可以得到对节点 j 发布的 j^0 类服务声明的各个因素的整体评价,评价值表示为六元组 $T = (j^0, T^A, T^B, T^C, T^D, T^E)$, 节点 j 声明的各项值越高,说明节点所声明的内容越符合实际,节点 j 也就越可信。这里存在着节点 j 的某一项声明的信任评价值很高,而另一项声明的信任值却很低的问题,通过六元组 T 的方法记录信任的评价值,能细化信任评价的粒度。

此时访问节点在选择服务时,不仅能够看到服务节点的声明,同时能看到对服务节点的声明信任评价,以下载速度为例,访问节点可以通过存放信任评价信息的档案点查看到 Web 服务的层次表及评价表,如表 3 所示。同理,也可以查看到声明 j^A , j^C , j^D , j^E 的相关列表。

表 3 声明 j^B 的 Web 服务层次表及评价表

层次	下载速度(j^B)	发布声明的服务节点	准确性评价
1	$j^B \geq 1\text{Mbps}$	$j_{m1}, j_{m2}, \dots$	$T_{m1}^B, T_{m2}^B, \dots$
2	$512\text{kbps} < j^B < 1\text{Mbps}$	$j_{n1}, j_{n2}, \dots$	$T_{n1}^B, T_{n2}^B, \dots$
3	$128\text{kbps} < j^B \leq 512\text{kbps}$	$j_{x1}, j_{x2}, \dots$	$T_{x1}^B, T_{x2}^B, \dots$
4	$64\text{kbps} < j^B \leq 128\text{kbps}$	$j_{y1}, j_{y2}, \dots$	$T_{y1}^B, T_{y2}^B, \dots$
5	$j^B \leq 64\text{kbps}$	$j_{z1}, j_{z2}, \dots$	$T_{z1}^B, T_{z2}^B, \dots$

通过表 3 可以看出,访问节点不仅可以看到提供服务节点的声明,而且也可以看到其它节点对该节点声明准确性的评价,这不仅有利于访问节点根据自身的网络带宽选择服务,而且能根据其它节点对提供服务节点声明的信任评价,选择高信任度的节点所提供的服务。

假设访问节点 i 是非常看重下载速度的节点,想得到 j^0 类服务,那么节点 i 首先通过信任评价信息的档案点查看例如表 3 所示的列表,然后根据自身的网络带宽,选择一些适合自己下载速度的节点 $j_x (x = 1, 2, \dots)$,再对这些节点进行综合信任评价,综合信任评价值 R_{ij} 表示为

$$R_{ij_x} = \mu \cdot WD_{ij_x} + (1 - \mu) (\overline{P_i^A} \overline{P_i^B} \overline{P_i^C} \overline{P_i^D} \overline{P_i^E}) \cdot (w_A w_B w_C w_D w_E)^T \quad (12)$$

其中 WD_{ij_x} 表示节点 i 对服务节点 j 的直接信任整体评价,通过式(6)获得; μ 表示信任评价的采用比例,这里取 $\mu = 2^{-(m_o/m_{ij_x})}$, m_{ij_x} 表示节点 i 与节点 j_x 直接交互的次数, m_o 表示节点是否有经验的门限值, μ 也可以通过用户自定义的方式获得。

通过式(12),节点 i 可以对适合自己需求的节点进行综合信任评价,采用综合信任值高的节点进行交互或合作,这样可以避免服务选择时的盲目性,能够提高服务选择的效率和交互的成功率。

2.3 分析

基于声明机制的信任模型,通过节点对自身的物理性能的声明,可以解决提供服务节点的物理差异问题;通过节点对服务的评价值采用六元组记录的方法,能够解决现有信任模型中访问节点评价标准差异的问题。该机制对节点发布不真实声明也有一定的抑制作用。

假设节点 j 是一个恶意节点,为了某种目的,故意发布不真实的声明。不真实声明分为两种:一种是夸大的声明,节点 j 发布的声明比自身的实际性能高;另一种是降低的声明,节点 j 发布的声明比自身的实际性能低。

对于夸大声明的节点,通过本文声明的准确性验证和信任评价,很容易就可以辨别,因为高性能节点更愿意跟自己性能相差不多的节点进行合作,对故意降低声明自己性能声明的节点,将失去跟高性能节点的合作机会,同时也可以通过建立激励机制的方法,促使节点发布真实的声明。

3 仿真实验

实验一:假设服务节点 j_1 是一个良好的节点, j_1 稳定地提供 j^0 类服务,数据包的下载速度为 50kbps,服务的响应时间为 0.2ms,提供服务的时间段为每天 7:00-23:00。节点 j_2 也是一个良好的节点, j_2 也稳定地提供 j^0 类服务,数据包的下载速度

为 20kbps,服务的响应时间为 0.4ms,提供服务的时间段为每天 7:00-20:00。节点 i 对节点 j_1 和 j_2 进行访问,并对这两个节点进行信任评价。

设 i 对节点 j_1 和 j_2 信任评价的初始值为 0.5,在不引入声明机制的情况下,可得到图 2 所示的信任评价;在基于声明机制的信任模型中,节点 i 对节点 j_1 和 j_2 的信任评价如图 3 所示。图 2 中,同一节点对不同性能的节点 j_1 、 j_2 的信任评价不同;图 3 中,同一个节点对不同性能的节点 j_1 、 j_2 的信任评价很相似,这是因为节点 j_1 和 j_2 所发布的声明跟实际行为相符。通过图 2 和图 3 的比较可以看出,在基于声明机制的信任模型中,只要节点发表的声明跟节点的实际行为相符,节点的信任评价就会很高。

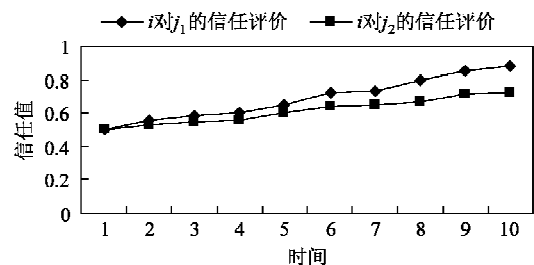


图 2 无声明机制时 i 对节点 j_1 和 j_2 的信任评价

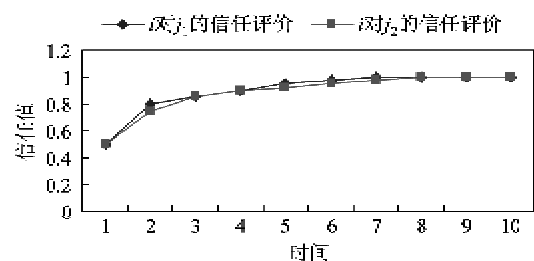


图 3 有声明机制时 i 对节点 j_1 和 j_2 的信任评价

实验二:假设服务节点 j 是一个稳定的节点,节点 j 在每天 24h 内稳定地提供 j^0 类服务,假设 j 所提供的服务安全性为 100%,所提供服务的品质是 80%,所提供服务的可靠性是 65%,并且在实验过程中这些数值保持不变。设节点 k_1 与节点 j 有丰富的直接交互经验,并对 j 所提供服务的品质评价为 100%,服务品质的评价是 80%,可靠性评价是 65%,节点 k_1 是一个非常看重服务可靠性的节点,节点 k_1 对 j 所提供服务的整体信任评价值是 70%,并且真实地向其他节点进行推荐。设访问节点 i 参考节点 k_1 对节点 j 信任评价,对 j 所提供的服务进行访问, i 根据自身的交互经验对节点 k_1 的反馈可信度进行更新,同时节点 i 非常看重的是服务质量,

并且 i 是一个善意节点, 提供真实的评价。节点 i 对节点 k_1 的反馈可信度进行更新(见图 4), 这里假设节点 i 对 k_1 的反馈可信度的初始值为 0.5。

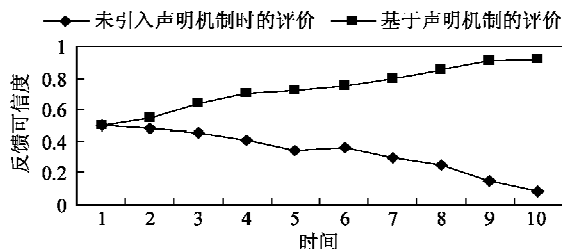


图 4 节点 i 对节点 k_1 的反馈可信度的更新

由图 4 可以看出, 在未引入声明机制的信任模型中, 对有经验的真实节点 k_1 的反馈可信度更新, 不仅没有增加反而降低, 这跟事实不符, 这种错误是因访问节点跟推荐节点对服务评价时评价的角度不同造成的。本文所提出的基于声明机制的信任模型, 可避免这种错误, 从而提高信任评价的准确度。

4 结 论

本文提出了基于声明机制的信任模型, 在一定程度上解决了分布式网络节点信任评价中节点的物理差异、评价标准差异及评价比较粗糙等问题。实验分析表明, 新模型精化了信任算法的粒度, 提高了

信任评价的准确度, 同时具有很好的包容性, 使得现有的信任模型可以对声明中的因素的准确性进行验证。

参考文献

- [1] Jøsang A. A logic for uncertain probabilities. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 2001, 9(3): 279-311
- [2] 徐锋, 吕建, 郑玮等. 一个软件服务协同中信任评估模型的设计. *软件学报*, 2003, 14(6): 1043-1051
- [3] Xiong L, Liu L. PeerTrust: supporting reputation-based trust for peer-to-peer electronic communities. *IEEE Trans TKDE*, 2004, 16(7): 843-857
- [4] 姜怡, 苏森, 陈俊亮. P2P 网络中基于实体行为的分布式信任模型. *高技术通讯*, 2005, 15(3): 1-4
- [5] 张骞, 张霞, 文学志等. Peer-to-Peer 环境下多粒度 Trust 模型构造. *软件学报*, 2006, 17(1): 96-107
- [6] 常俊胜, 王怀民, 尹刚. DyTrust: 一种 P2P 系统中基于时间帧的动态信任模型. *计算机学报*, 2006, 29(8): 1301-1307
- [7] 林齐宁. 决策分析. 北京: 北京邮电大学出版社, 2003
- [8] 汪应络. 系统工程. 北京: 机械工业出版社, 2003
- [9] Ratnasamy S, Shenker S, Stoica I. Routing algorithms for DHTs: Some open questions. In: *Proceedings of the First International Workshop on Peer-to-Peer Systems*, Cambridge, USA, 2002. 45-52

A Web trust model based on declaration mechanism

Deng Zhongjun*, Wang Shaojie**, Zheng Xuefeng*, Lin Ran***, Suo Yanfeng*, Yu Zhen*

(* School of Information Engineering, University of Science and Technology Beijing, Beijing 100083)

(** National Research Center for Information Technology Security, Beijing 100084)

Abstract

The paper proposes a new trust model based on declaration mechanism to solve the evaluation differential problems in peers, trust evaluation in Web service, including the aspects of the physical differential of peers' capabilities in connection and computation, the subjective differential of peers' preferences in trust evaluation, and the comprehensiveness differential on trust evaluation factors. The trust model takes the physical differential and subjective differential among peers into consideration, and sets the evaluation criterion of trust in detail. The experimental results show that the proposed trust model can thin the granularity of the trust evaluation criterion and improve the accuracy of trust evaluation. Furthermore, the model also embodies the personal characteristic of peers and better scalability.

Key words: Web service, declaration mechanism, trust