

基于滑动窗口的云服务可信评估模型优化^①



李 圳, 于学军

(北京工业大学 信息学部, 北京 100124)

通讯作者: 李 圳, E-mail: 18519020056@139.com

摘 要: 随着云服务的需求日益增强, 云服务市场得到了快速的发展, 但也产生了云服务的可信问题. 国内外针对这个问题提出了各类的云服务可信评估模型, 使用直接间接可信, 用户评价、偏好相似度等因素参与云服务可信评估的流程中, 然而现有研究中的云服务可信评估模型在流程中针对信任评估中并不能满足像是缓慢增加和快速下降等基本性质. 本文引入云服务可信滑动窗口的方式加强整个云服务可信评估的流程, 通过滑动窗口满足了缓慢增长等性质解决了问题, 而后通过模拟实验验证了确实能够有效解决该问题.

关键词: 云服务可信性评估; 滑动窗口; 言行一致

引用格式: 李圳, 于学军. 基于滑动窗口的云服务可信评估模型优化. 计算机系统应用, 2019, 28(7): 35-43. <http://www.c-s-a.org.cn/1003-3254/6984.html>

Optimization of Cloud Service Trusted Evaluation Model Based on Sliding Window

LI Zhen, YU Xue-Jun

(Faculty of Information Technology, Beijing University of Technology, Beijing 100124, China)

Abstract: Along with the increasing demand for cloud services, the cloud service market has developed rapidly, but it also produced a trustworthy problem with cloud services. Various cloud service trust evaluation model have been proposed for this problem domestically and internationally. They Use direct trust and indirect trust, user evaluation similarity and preference similarity to participate in the process of cloud service trustworthy assessment. However, the cloud service trustworthy evaluation model in the existing research cannot meet the basic properties such as slow growth and rapid decline in the process of trust evaluation. This paper introduces the way of cloud service trustworthy sliding window to strengthen the process of trustworthy evaluation of the entire cloud service. The problem was solved by using a sliding window to satisfy the slow growth property, and then it was verified by simulation experiments that it can effectively solve the problem.

Key words: cloud service credibility assessment; sliding window; consistent

在信息化时代的 21 世纪, 软件作为信息的载体之一, 已经广泛应用到经济、军事以及日常生活等各个方面, 成为人类和社会发展不可或缺的工具之一. 然而, 随着软件规模的不断扩大, 其内部结构越来越复杂, 应用环境越来越开放, 软件的不可信给人类社会带来的各种不安全因素也随之产生. 因此, 提高软件可信性具

有重要意义.

一般认为, “可信”是指一个实体在实现既定目标的过程中, 行为及结果可以预期, 它强调目标与实现相符, 强调行为和结果的可预测性和可控制性, 软件的“可信”是指软件系统的动态行为及其结果总是符合人们的预期, 在受到干扰时仍能提供连续的服务, 这里的“干扰”

① 基金项目: 国家重点研发计划 (2017YFF0211801)

Foundation item: National Key Research and Development Program of China (2017YFF0211801)

收稿时间: 2019-01-17; 修改时间: 2019-02-03; 采用时间: 2019-02-19; csa 在线出版时间: 2019-07-01

包括操作错误、环境影响、外部攻击等^[1].而软件行为声明,就是基于“言行一致”思想,站在软件全生命周期视角,提出基于行为声明的软件可信性测试的解决方案,从而达到最终完成的软件产品具有相对较高的可信性.

经研究发现目前云计算商业领域除了有云平台厂商自身提供的一些大数据处理类的服务以及服务器租借服务外,还提供了含有第三方云服务提供商的云市场模块,目前有云市场功能的云平台(主要是用友云和阿里云)都并未使用任何有效的措施约束第三方云服务提供商提供的服务质量,而是简单的靠着在开发商手册中的条款希望云服务提供商能够自觉遵守协议,服务的反馈机制仅仅是靠着用户的打分和一点评论,很多用户并不在意这种评论而乱评论,是一种被动可信评估的机制,这种机制明显不能解决云服务商提供服务的可信问题,所以目前急需一个完备的云服务可信性评估与不可信服务筛选的机制.

目前国外与国内均对于服务可信相关问题有所研究,国外更多着眼于信息可信度相关的研究,大多数研究都参考了 Josang^[2]有关于用户信誉度评估相关的论文,更多是分析服务用户可信,如 Shojaee 等人^[3]提出的一种基于用户的正确和恶意反馈来计算出用户信誉度的方法, Peng 等人^[4]提出了一种解决新进用户初始化信誉度的方法,根据用户行为识别和分类来初始化用户信誉度, Noorian 等人^[5]针对用户信誉提出了一种推荐人机制,并提出相应的鼓励和惩罚措施, Vassileva 等人^[6]又根据用户的习惯倾向性,喜好区别以及经验丰富程度等因素对用户信誉计算进行了下一步的优化,除了国外研究者,国内也有杨蕊岚等人^[7]使用多级模糊评估模型评估云计算中用户可信度的研究.国内更多偏向服务端的可信,主要基于路径的服务可信研究,主要包括集群相关,偏好相关的服务可信评估,最优服务组合选择等.曹洁等人^[8]为了解决数据的不确定性提出了利用信息熵理论的评估权重算法.王海艳等人^[9]为了解决用户间信任计算的问题,将评价数据相似的用户分为集群计算.盛国军等人^[10]为了优化选择最可信调用路径使用了蚁群算法.虽然现今国内外对于云服务的可信评估都有所研究并建立了各自的信任评估模型,但是至今为止提出的各模型都并不能满足可信评估模型的一些基本性质如信任值缓慢增长、信任值快速下降等.

为了解决现有研究中云服务信任评估模型不能应对恶意云服务提供商的信任滥用和重入等恶意为对

信任评估准确性的影响,本文引入滑动窗口机制优化云服务信任评估流程,滑动窗口通过双窗口的结构以及为信任评估流程带来信任缓慢积累以及快速下降等性质.之后提出了基于滑动窗口的云服务可信评估模型,并通过实验验证了信任评估模型的有效性.

1 云服务可信评估

1.1 云服务可信

信任最开始从人与人之间的信任得来^[11],本来是一方对另一方是否能守约的主观感觉,所以更多的是要强调“言行一致”,即事先声明(约定)与实际的结果相一致.在云服务中信任主要代表云服务方事先声明的服务属性和实际使用时的服务属性是否一致.

云服务的“言行一致”则是云服务提供商在平台上注册服务时声明的服务质量与实际用户在使用时的服务质量是否一致来评定,每次用户在评定服务是否可信后将反馈数据传回平台,而平台将会根据这些数据帮助之后的用户筛选出可信的服务.

1.2 云服务可信评估的方法

图1为云服务可信评估模型的全貌,主要评估因素就是该云服务历史使用时的反馈记录,而反馈记录由反馈来源的不同分为了用户自身对于目标服务的反馈数据和其他用户对于目标服务的反馈数据,这两种反馈数据计算后分别得到目标服务的直接信任评估值和间接信任评估值,最后将两部分的信任评估结果整合后得出目标服务的可信评估值,下面将对流程中每一个组成部分进行介绍.

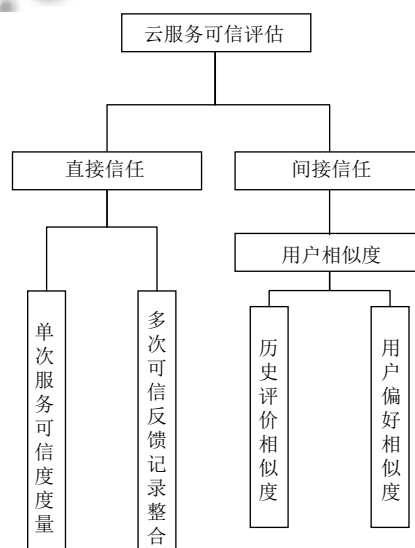


图1 云服务可信评估方法示意图

1.2.1 直接信任

直接信任,即为用户自己在以前与服务交互的记录,直接信任由于是用户自己本身的反馈,假定条件为不可能存在造假或是欺骗的情形^[12],由于对于用户而言,自身是完全可信的,是一种主观信任,所以应在整个评估环节占主导地位.直接信任值主要由用户自身的多次反馈数据所综合得出.主要涉及到了单次服务信任值的计算和带时间戳的多次信任评估值的综合数值计算.

(1) 单次服务信任值

单次信任评估值,即用户在每次使用完服务后的反馈数据评估值,是整个评估模型中存储数据的基本单位,充分贯彻了“言行一致”的理念,为服务提供商在注册服务时声明的服务质量和实际用户在使用时反馈值的对比数值^[13],从这个比值中能得出单次调用时该云服务是否可信.同时还要注意标准化问题^[14].

(2) 多次信任评估值的综合数值

平台根据用户对于目标用户的多次信任反馈数据根据相应算法得出综合信任值即直接信任.

1.2.2 间接信任

间接信任^[15]代表的是其他用户对目标服务的信任数据.间接信任由于并不是当前用户自身的反馈数据,所以除了要计算前面 1.2.1 节中提到的多次反馈数据总计合计算出的信任值外,还需要计算用户间的相似度.相似度主要分为用户偏好相似度和用户评价相似度.

(1) 用户偏好相似度

正如 1.2.1 节直接信任中单次反馈信任值计算中叙述,用户在反馈实际服务质量数据到算出信任值时不同的服务质量属性可以设定不同的偏好值,而用户间偏好分配不同可以导致对于同一个服务同样的实际服务质量产生不同的信任值,所以间接信任值计算中需要计算用户偏好相似度.

(2) 用户评价相似度

不同的用户除了在计算时偏好不同会有所影响外,使用服务的时段也有所不同,所以可以根据用户评价来筛选选用服务更相似时间也相似的用户^[16].同时可以抵消而已云服务共谋行为对结果带来的影响.

本节介绍了云服务信任的定义以及现有研究中云服务评估模型的结构,虽然现有研究中的评估模型使用多种机制来评估云服务是否可信,但是现有的机制并不能解决诸如恶意云服务重入或信任滥用等问题,

接下来本文将会提出云服务信任评估应满足的基本性质,并通过引入滑动窗口优化云服务信任评估流程.

2 基于滑动窗口的云服务可信评估流程

2.1 基于滑动窗口云服务评估模型因素分析

对应云服务信任的评估中,对于同一个用户对同一个目标服务的多次反馈记录进行融合出直接信任是估算云服务可信评估值的第一步,而为了能够更好的筛选服务,如缓慢增长和迅速下降两个最基本的性质要得到满足,而国内外对云服务可信评估中这一环大多采用单时间衰减因子来参与计算,但是单衰减因子并不能解决常见的恶意行为,如重入行为,被评估为不可信的服务通过改变名称重新注册服务来清楚过去的反馈记录,不可信的云服务提供商还会在注册初期积累可信记录后提供不可信记录仍然还能被评估为可信云服务.现有的云服务信任评估模型并不能很好的解决上述恶意行为对评估流程带来的影响导致整个评估模型对于不可信云服务的鉴别效率低,此外还有一些性质也不能得到很好的满足,而通过本文加入的滑动窗口方法能较好地解决性质无法得到满足的问题,这些性质分别为缓慢增长、快速减少、事件相关性、记录有效数量以及奖惩制度,下面对 5 个基本性质进行叙述.

性质 1: 缓慢增长

现实中对于商品的口碑是要慢慢积累的,少量的好评并不能代表商品质量就是优异的,云服务同样,几次零星的可信调用并不能直接判断当前服务可信,如正向积累值的增长速度要慢,另外通过极少数评价来断定是否可信会导致不可信的云服务商可以以极小的成本积累一些少量的可信反馈从而达到被评估为可信的目的.

性质 2: 快速减少

对于云服务产生的负面值反馈必须比正向的增长得到更多的影响,及一旦出现不可信的调用反馈立即对于信任预测值进行大量的下降处理,以作惩罚作用.

性质 3: 时间相关性

用户对于信任值的预测必须具有时间相关性,信任评估值应该反映的是云服务近期时间段中的可信表现,而随着服务可信表现的波动服务的信任评估值也要随着做出改变,同时随着长时间没有使用记录,服务的可信评估值应该慢慢趋近于初始值,综上所述需要找到

一个方法使评估时新的记录所占权重比旧的记录要高。

性质 4: 记录有效数量

随着使用的次数渐渐积累, 对于评估的消耗会越来越大, 而一些很久以前的记录对于最终评估结果能产生的影响微乎其微而占用同样的计算时间, 所以对于过多的累计记录应该有一定的有效记录数量的限制。

性质 5: 服务可信评估值的奖惩制度

当有多个服务可信反馈评估值记录时, 对于评估值不同的反馈 (高于信任阈值及可信的反馈值和低于信任阈值及不可信的反馈值) 在评估中应对最终计算出的直接信任值有不同的影响, 可信记录应该使最终的直接信任值增加, 而不可信记录应该使最终的直接信任值降低, 对应云服务能够和不能够提供可信服务时对于最终评估值产生的影响。

为了满足以上 5 个性质, 主要使用两个措施, 分别为时间衰减因子和滑动窗口。时间衰减因子作为云服务可信数据整合中常见的步骤, 代表了不同时间的反馈值在整合中占有了不同的权值, 确保离当前评估时间越近的反馈有更高的权值。时间衰减因子满足了时间波动性, 因为引入了时间衰减因子后离当前评估时间越近的记录所占比重越大。而其他的几个性质中性质 2 和性质 5 的效率很低, 只是单纯使用了时间衰减因子的权重下不是很好的能够在短时间内筛选出不可信的服务, 尤其是该服务如果长时间提供的是可信的服务时服务质量产生了波动从而不可信后评估的响应速度过低, 而性质 1 和性质 4 则是完全没有做到, 所以现今常用的时间衰减因子是需要进行优化的。

为此本文使用滑动窗口的方法, 滑动窗口通过填充窗口数值, 不可信惩罚措施以及限制窗口大小等措施来满足现有研究中在效率上不能达到的问题。

2.2 基于滑动窗口云服务评估方法

作者引入滑动窗口解决 2.1 节提到的目前云服务信任领域单纯依靠时间衰减因子处理反馈数据整合所引发的服务可信评估值的奖惩制度和快速减少不能很好的达成以及缓慢增长和记录有效数量完全不能达成的问题。具体设计如图 2 所示。

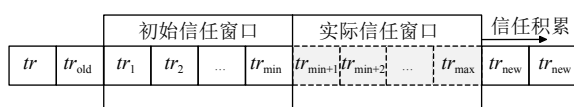


图 2 云服务可信评估所用滑动窗口示意图

如图 2 中所示, 云服务可信评估所用滑动窗口主要是两部分, 初始信任窗口和实际有效窗口^[17], 图中每一列 tr 代表信任反馈数据, 后缀中 $1-\min$ 为缓慢增长窗口, $\min+1-\max$ 为实际有效窗口的记录, 其中新的反馈数据从右侧作为 tr_{new} 进入窗口中, 而超过窗口大小时最久远的数据作为 tr_{old} 从窗口左侧淘汰, 整体窗口向右行进。通过其中的特性能够解决前面所说现有研究中无法解决的有效记录、奖惩制度, 缓慢增加, 快速减少的问题。下面将对滑动窗口的主要组成部分和所解决的问题进行详细说明。

(1) 窗口限制

首先滑动窗口使用了普通网络传输中窗口中类似的功能, 将反馈数据存储在窗口中进行评估, 这样使用窗口大小限制了有效记录的数量, 从而解决前面记录有效数量的问题。

(2) 窗口分块之信任初始窗口

除了使用窗口大小限制有效次数外, 滑动窗口为了满足其他性质还将整个窗口分为两个部分, 分别为信任初始窗口和实际有效窗口。

信任初始窗口用来在历史反馈数量不太多的时候让信任以缓慢增长的方式来变化, 除了自带时间衰减因子外, 通过将初始值填满整个信任初始窗口并且时间自动设置为评估时间以减小在信任初始窗口中可信反馈记录对信任评估值的增加效果, 以达到缓慢增长的作用。

(3) 窗口分块之实际有效窗口

当信任反馈记录超出信任初始窗口后记录开始进入实际有效窗口中, 与初始窗口不同, 有效窗口并不会使用初始值填充窗口, 而是直接使用实际数据, 而且限制了有效记录次数, 当次数超过实际窗口最大值时新来的记录仍然填进实际窗口中, 而从信任初始窗口弹出最老的数据。

(4) 滑动窗口的奖惩措施

前面的功能已经解决了大部分问题, 就剩下惩罚措施和快速下降, 其中快速下降就是惩罚措施的主要表现手段, 及之前云服务提供商在一开始一段时间内提供的是可信的服务, 然而突然有一两次有了不可信的反馈数据则可信评估值迅速向低于可信阈值的方向转变。滑动窗口的惩罚措施主要是对出现恶意反馈之前高于可信阈值的记录, 将这些反馈数据改变为初始化记录, 这样通过将之前的可信反馈数据修改为初始

值,从而使云服务信任评估值快速下降,也就实现了惩罚的作用。

通过以上叙述的几个措施,云服务信任滑动窗口通过两个不同且限制大小的窗口以及特别的惩罚措施解决了现有云服务评估研究中未能解决的问题,从而提高了信任评估的准确性以及在持续使用中产生波动时的调整效率。

2.3 评估流程与优化

前面几节介绍了评估模型的组成部分以及现在模型中存在的问题以及可以优化的点,下面对云服务可信评估整个流程进行详细介绍。首先整个流程示意图如图3所示。

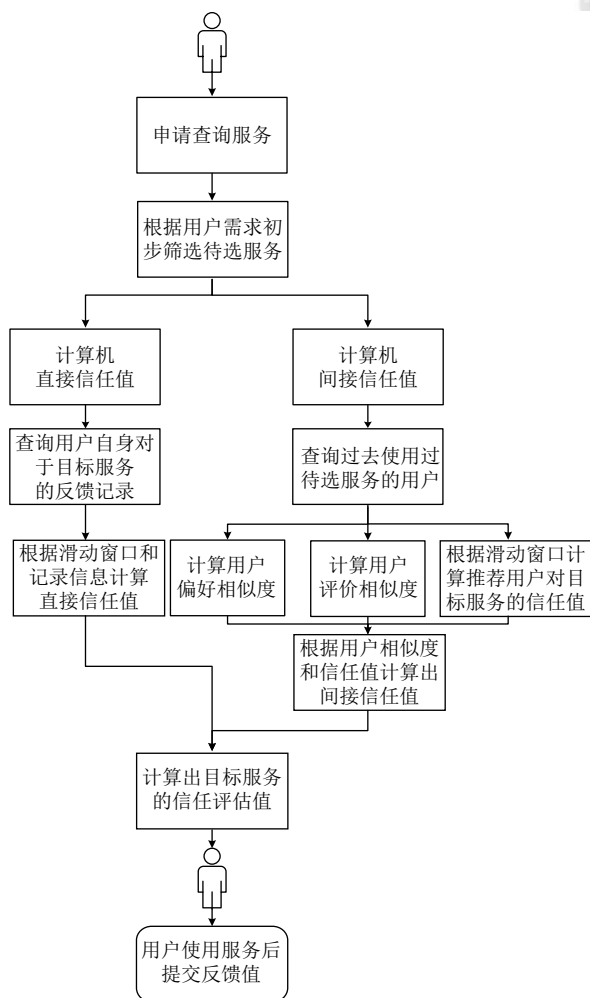


图3 云服务可信评估流程示意图

整个流程中主要有三个主要角色,分别为用户,服务提供商,以及平台。

平台在整个实验过程中起到转交数据和根据数据

评估可信值,是其中最重要的一环。

用户负责提交需求和根据筛选后的服务列表选择服务。

服务提供商负责根据自己在设定中的角色生成对应的服务数值。

下面分别详细介绍流程图中每一步的细节以及所用公式。

(1) 每次调用时用户选取一定服务质量的需求提交给服务器。

评估的前提便是用户有使用服务的需求,接下来交给平台对服务进行筛选。

(2) 服务器以用户需求为最低限度的要求搜索服务质量大于该值的服务,完成初步筛选。

首先服务的最基本要求是符合用户需求,首先根据用户搜索服务的类别和属性需求对服务进行初步筛选。

(3) 平台根据用户自身对于目标服务的反馈数据计算直接信任值。

直接信任的评估是根据滑动窗口共识计算得出的,这样做能够解决章节2.1中提出的问题并优化流程。

当次数少于缓慢增长窗口大小时公式如下:

$$T_{\text{straight}} = \frac{\sum_{k=1}^n T_k * e^{\wedge(t_{\text{rec}} - t_{\text{curr}})} + (n_{\text{min}} - n) * T_{\text{new}}}{\sum_{k=1}^n e^{\wedge(t_{\text{rec}} - t_{\text{curr}})} + (n_{\text{min}} - n)} \quad (1)$$

其中, n 为记录次数 n_{min} 为缓慢增长窗口大小, t_{curr} 为当前时间, t_{rec} 为记录时间, T_k 为单次信任评估值, T_{new} 为信任初始值。

当次数大于缓慢增长窗口时公式如下:

$$T_{\text{straight}} = \frac{\sum_{k=1}^n T_k * e^{\wedge(t_{\text{rec}} - t_{\text{curr}})}}{\sum_{k=1}^n e^{\wedge(t_{\text{rec}} - t_{\text{curr}})}} \quad (2)$$

其中, T_{straight} 是直接信任值, t_{curr} 为当前时间, t_{rec} 为记录时间。

(4) 平台计算间接信任值。

平台根据过去与目标服务有交互且和当前用户有共同使用的用户所有的历史反馈数据计算间接信任值

(5) 服务器从数据库中查询每一个服务过去的用户反馈记录。

首先要选择推荐用户,推荐用户即为非当前请求服务用户外过去使用过目标服务的用户,查询这些用

户过去使用目标服务的记录。

(6) 服务器根据其他用户和当前请求用户对于服务质量属性权重的区别计算出用户相似度。

用户相似度主要根据用户偏好相似度 (*Simp*) 和用户评价相似度 (*Sime*) 来得出, 下面对这两个属性分别做介绍。

① 用户偏好相似度

用户偏好相似度主要代表的是用户对于各 QOS 属性在评估时设定的权重不同所带来的差别, 当用户间偏好较大时应该在间接信任评估过程中所占权重要低, 权重所用公式如下:

$$Simp = \frac{\sum_{k=1}^n |w_{1k} - \bar{w}_1| |w_{2k} - \bar{w}_2|}{\sqrt{\sum_{k=1}^n [w_{1k} - \bar{w}_1]^2} \sqrt{\sum_{k=1}^n [w_{2k} - \bar{w}_2]^2}} \quad (3)$$

其中, *Simp* 就是计算出的用户偏好相似度, *n* 为 QOS 属性的个数, w_{1k} 为当前请求用户对每个属性的偏好权重, w_{2k} 为待选推荐用户对每个属性的偏好权重, 此公式采用的是皮尔森系数公式, 为常用的相似度计算公式。通过这个相似度计算公式能够反映出不同的用户调用同一个云服务在获取到的服务质量属性相同的时候由于偏好不同所产生的差别, 提升间接信任评估的准确性。

② 用户评价相似度

除了使用偏好相似度反映用户偏好所带来的差别外, 用户间非目标服务的评价也会有所差别, 用户的使用服务次数, 时间段, 以及可以通过一定手段为商家服务的恶意用户, 都可以通过用户评价相似度有效地将这些问题的影响减少到最小。其中用户评价相似度首先通过目标用户和请求用户间共通的云服务使用记录为条件筛选推荐用户, 没有共通服务使用记录的首先就会被排除, 之后通过每个共通服务的欧氏距离公式计算出相似度, 公式如下:

$$Sime = \frac{\sum_{j=1}^m \sqrt{\sum_{k=1}^i (T_{j,k} - T_{j,k'})^2}}{m} \quad (4)$$

其中, *Sime* 为计算出的用户评价相似度, *i* 为 QOS 属性数量, *m* 为请求用户和目标推荐用户共通的服务数量, $T_{j,k}$ 为用户对目标云服务的直接信任反馈值。

将上面的用户评价相似度和用户偏好相似度整合起来就会得到用户相似度, 公式如下:

$$Sim = Simp / Sime \quad (5)$$

其中, *Sim* 代表最终计算出的用户相似度, *Simp* 为用户偏好相似度, *Sime* 为用户评价相似度, 由于用户偏好相似度使用的是皮尔森系数公式, 在用户间偏好完全相同时为最大值 1, 偏好差距越大得到的值最小, 而用户评价相似度使用的是欧几里得距离公式, 所以是越相似值越小, 所以用户相似度为用户偏好相似度和用户评价相似度的比值。

(7) 根据用户相似度, 时间窗口和每一个用户反馈的可信值计算出间接信任值。

间接相似度在计算出了用户相似度外就要计算推荐用户对目标云服务的直接信任, 而直接信任即是本模型与之前研究最大的一点不同, 通过滑动窗口解决奖惩措施、缓慢增长、快速降低、有效次数的问题, 从而提升云服务可信评估的有效性。

除此以外还有限定总窗口大小以及惩罚更新反馈数据的步骤, 这些一起组成了云服务信任评估模型中的各个措施, 增强了评估的准确度。

(8) 使用请求用户自己反馈的可信值以及间接信任值计算出总的可信评估值。

使用用户自身的直接可信以及前面计算出的间接可信, 最终公式如下:

$$T = T_d * w_d + w_{in} \sum_{k=1}^y sim_k * T_k \quad (6)$$

其中, *y* 为推荐用户数量, T_d 为直接信任值, w_d 为直接信任权重, w_{in} 为间接信任权重, sim_k 为用户相似度, T_k 为推荐用户的直接信任值, 最终得出的便是云服务信任评估值。

(9) 最终筛选完的服务既符合可信阈值又满足客户对于服务质量的要求, 并转交给用户。

(10) 用户从列表中选择一个服务并使用该服务。

(11) 平台获悉用户选择了服务后通知服务提供商提供服务。

(12) 服务提供商根据自己在数据库中的服务质量和自己的角色 (恶意, 正常, 偏向) 生成实际服务质量传回给平台。

(13) 用户调用服务结束后反馈 QOS 数据, 平台根据用户的偏好属性计算出单次使用的云服务信任评估值, 并存储。

其中平台需要读取用户对各 QOS 属性的偏好, 读

取云服务事先声明的属性并进行计算,公式如下:

$$Ts = w_1 * \frac{s_1}{q_1} + w_2 * \frac{s_2}{q_2} + \dots + w_n * \frac{s_n}{q_n} \quad (7)$$

其中, n 代表 QOS 属性数量, $w(1-n)$ 代表当前用户对于 QOS 属性的偏好权重, $s(1-n)$ 代表实际调用时的 QOS 属性值, $q(1-n)$ 代表云服务商在注册时生命的 QOS 属性值 (SLA)。

计算出 Ts (单次云服务信任评估值) 后平台将会把数据进行存储, 单次云服务信任评估值作为云服务可信评估模型的基础组成部分, 接下来进行的云服务信任评估也要基于这些数据进行计算。

(14) 平台根据世纪服务质量计算实际信任值并存储进数据库。

以上便是用户正常调用时云服务可信评估模型在其中执行的流程, 通过在整个评估过程插入滑动窗口来优化整个模型, 下面将通过一个模拟实验来验证整个模型的优越性。

3 实验仿真及结果分析

为了验证基于滑动窗口的云服务可信评估模型优于当前研究中常用的直接间接信任模型, 现在对带滑动窗口和不带滑动窗口的云服务可信模型进行对比实验, 实验所用环境为 Win10 系统, 16 GB 内存, 所用模拟实验的架构为 SpringBoot+Mybatis+Redis。

为了尽可能地模拟云服务环境下的情况, 模拟实验中定义三个角色, 分别为用户, 服务提供商, 以及平台。平台在整个实验过程中起到转交数据和根据数据评估可信值, 是其中最重要的一环, 用户负责提交需求和根据筛选后的服务列表选择服务, 服务提供商负责根据自己在设定中的角色生成对应的服务数值。

本次实验因为是要验证研究的优越性, 即在突出了间接信任中常用的偏好相似度和评价相似度以及使用到常用的时间衰减因子的情况下对比可信评估成功次数

在直接信任值计算中, 作者选取了滑动窗口的是否使用来进行对比试验, 带时间窗口的实验使用了时间窗口和时间衰减因子综合得到的结果用来在实验中计算直接信任值, 而在对比试验中使用了单纯的时间衰减因子来进行运算。

为了模拟出时间窗口所需要解决的情景, 对于云服务提供商的恶意行为, 实验选用了信任滥用和重用两种常用恶意行为, 在实验中设定了第二种服务商, 即

为恶意云服务提供商。恶意服务商在前半段时间中积累可信记录, 以重入的方式快速积累正向信任值, 之后利用之前积累的信任开始提供大量的不可信服务, 从而达到为自身牟利的目的。该恶意云服务主要包括两种状态, 信任积累状态与恶意滥用状态, 实验中平台使用的信任阈值为 0.8, 当单次调用的信任反馈值或信任评估值超过 0.8 即认为该服务可信。信任积累状态恶意云服务提供商和普通云服务提供商提供的服务总值根据自身声明的服务质量与随机函数提供 0.9–1.0 信任值的服务, 在信任滥用阶段云服务提供商提供的服务总值根据自身声明的服务质量与随机函数提供 0.65–0.85 信任值的服务来模拟信任滥用恶意行为。对于用户侧的恶意行为, 本实验选用马甲诋毁与马甲共谋两种行为, 正常用户在使用完服务后根据实际 QOS 属性数据和云服务事先声明的属性值计算比值后根据自身的偏好权重整合后得到单次服务的信任值后上传, 而恶意用户在与共谋的恶意云服务提供商改变为信任滥用模式后, 当其调用正常的云服务时执行马甲诋毁行为反馈低于信任阈值 (即判定为不可信) 的反馈值, 而对于与恶意服务则执行马甲共谋行为, 反馈高于可信阈值的反馈值。通过在流程中分别添加恶意云服务提供商与恶意云服务用户来模拟实际使用时恶意行为对云服务信任评估流程带来的破坏效果。最后通过每轮的可信服务 (服务中实际的属性数据/服务厂商声明的服务质量 $\geq$ 服务阈值) 占有所有进行的服务的比值来观察带时间窗口和不带时间窗口时服务服务可信率的回弹速度来显示时间窗口的优越性。

除了需要模拟恶意云服务提供商以及与其共谋的恶意用户对云服务信任评估流程的破坏行为外, 也需要考虑普通用户在需求不同时对于云服务可信评估准确性的影响, 为此除了用户在每次调用时会使用随机的需求以及云服务设置了不同的云服务声明 QOS 外, 也需要模拟用户对于不同偏好下对于评估结果造成的影响。由于用户对不同的于对比带不带偏好而产生的相应区别, 服务提供商的构成中多加了一种“偏好服务提供商”, 该提供商对于其中某一属性的服务质量很高, 只能对该属性偏好下的用户提供服务, 在本实验中该偏好服务提供商提供的服务为响应时间偏好服务, 在响应时间中提供的是 0.9–1.0 倍声明的属性, 而对于吞吐量则提供的是 0.6–0.65 倍声明属性, 剩余两个 QOS 属性提供的是 0.8–0.9 倍, 这种类型的云服务虽然在吞吐量偏好的用户进行反馈时会被判定为不可信,

但是对于响应时间偏好用户来看是可信的,从而云服务评估系统应该有能力将响应时间偏好的服务推荐给响应时间偏好用户,而不是根据吞吐量偏好用户低于信任阈值的反馈记录就判定该服务为不可信。同时在用户定义时将用户分为两类:响应时间偏好用户和吞吐量偏好用户,响应时间偏好用户各服务质量偏好权重为可用性:可靠性:响应时间:吞吐量=2:2:5:1,吞吐量偏好用户各服务质量偏好权重为可用性:可靠性:响应时间:吞吐量=2:2:1:5。此时若响应时间偏好云服务生成的实际属性分别为0.8,0.8,0.9,0.6时,在响应时间偏好用户的权重分配下计算出的反馈值为0.83,超过可信阈值0.8,而如果是吞吐量偏好用户的权重分配下计算出的反馈值为0.71,低于可信阈值,可以看到同样的云服务质量属性下在不同偏好用户的评判下会产生不同的结果。此实验主要强调的是偏好相似度会加强相似偏好的用户的反馈在全部用户反馈中的分量,从而将相同属性的偏好服务器推荐给该属性偏好的用户,而不会推荐给非该属性偏好的用户。同样,偏好服务会在前一半时间只为同偏好用户提供服务来积攒偏好反馈数据。

实验使用多线程模拟的形式进行,使用线程池并行调用的方式创建多个用户对象调用服务并进行轮次调用,在使用时调用云服务信任筛选服务后选用评估为可信的云服务中其一调用并计算当前调用的信任反馈值并记录,由于记录为单用户对目标服务的线性记录所以针对数据结构进行了一些优化。实验中创建了180的用户对象,其中包含60个响应时间偏好用户,60个吞吐量偏好用户以及60个恶意用户,选用的服务一共有150个,其中包含50个恶意云服务,50个普通云服务,50个吞吐量偏向云服务,由于为了对比云服务可信评估模型对于各类恶意行为造成的破坏效果的修复和屏蔽效果,本实验设置了行为转变信号,指一开始时恶意用户和恶意云服务都进行正常的调用和服务,云服务在此时积攒信任,之后进行到一定轮次后恶意云服务开始利用自身积攒的信任记录提供不可信的云服务,而恶意用户开始进行诋毁和共谋操作,每轮每个用户请求固定数量的云服务并调用后并反馈,在所有用户调用完固定次数后该轮结束,平台根据当前轮云服务调用反馈中超过可信阈值的调用(可信调用)次数和当前轮全部用户的总调用次数的比例即为当前轮次中可信调用所占比例。

图4为实验结果。

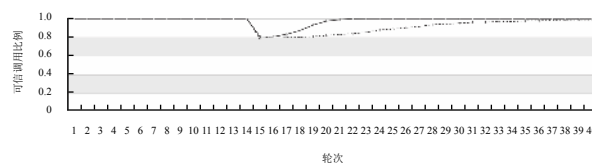


图4 滑动窗口对比实验结果

图4中横坐标为可信评估的轮次,其中每一轮有180名用户分别进行15次调用,一共2700次服务调用,总共40轮108000次调用数据,纵坐标是每轮中经过服务可信评估筛选后的调用中最终可信调用(即服务上提供了言行一致的云服务)的比例。从图中可以看出在第14轮后两个实验的可信服务成功率都有很大程度的下滑,然而带时间窗口那一次(实线)明显比不带时间窗口(虚线)更快的将成功率提升回之前的接近1的水平,正是因为带时间窗口时服务的可信值增长幅度较低,而且会在产生非可信服务后快速降低可信值,从而更加迅速的筛选掉了恶意服务提供商。

上面的实验结果证实了在计算直接信任值时时间窗口的重要性。

服务中使用的数值如下:

时间衰减因子:使用的是 $1.5^{-(\text{实验轮数差})}$,同轮下权值最大,为1。

时间窗口中缓慢增长窗口大小:50

可信服务阈值:设定为0.8,低于0.8的可信值被定义为不可信。

恶意服务提供商变更行为模式信号:14,也就是说14轮后开始提供非可信服务。

偏向服务开始对全部用户开放的时间:14

不同偏向用户对于属性的权重值:0.2,0.2,0.5,0.1和0.2,0.2,0.1,0.5

这样不同类别用户对于同一服务的评价会有差别

恶意服务提供商变更行为模式信号:14,也就是说14轮后开始提供非可信服务。

恶意用户行为模式:恶意用户在14轮之前行为和普通用户并无区别,在14轮后除了正常服务返回正常的反馈外,对恶意服务商默认返回可信反馈(0.95)。

由于基于滑动窗口的云服务评估模型中使用滑动窗口对流程进行优化,相比现有的模型在本文提出的模型中恶意服务在积累信任时的速率较慢,则云服务商通过短期正常行为获得的可信度比原有的更低,而当其结束重入后开始信任滥用行为时其信任快速下降至低于信任阈值,从而云服务的调用成功率相比现有

的模型更快速回到了正常值。

4 结束语

为了应对云计算环境下云服务产生的可信问题,国内外的研究者提出了各自的云服务评估模型以评估云服务是否可信。本文在对云服务信任评估相关的研究进行了较为全面的总结与分析的基础上,对云服务信任评估流程进行了一定的优化与改进,构建了基于滑动窗口的云服务信任评估模型,本文的研究可以总结为以下几个方面:

(1) 总结现有研究中的云服务信任评估模型并分析出模型中待完善和改进的方面,基于常见恶意行为以及其他评估流程中需要关注的需求提出云服务信任评估模型应满足的基本性质。(2) 基于云服务信任评估的基本性质引入滑动窗口对现有云服务信任评估流程进行优化和改进,使其满足基本性质。(3) 根据现有模型以及本文引入的滑动窗口,构建基于滑动窗口的云服务信任评估模型。同时通过第3节的仿真实验可以看出,本文提出的基于滑动窗口的云服务信任评估模型相比现有的评估模型能够更快的筛除不可信的云服务。

本论文在引入了滑动窗口优化云服务可信评估时,只考虑了按单个服务情况下的云服务信任评估,而在一些复杂场景是需要使用多个服务的,所以未来可以对服务组合情况下的云服务可信评估进行进一步研究,通过评估来选出最可信的云服务组合。

参考文献

- 1 刘克,单志广,王戟,等.“可信软件基础研究”重大研究计划综述. 中国科学基金, 2008, 22(3): 145-151. [doi: 10.3969/j.issn.1000-8217.2008.03.005]
- 2 Jøsang A, Ismail R. The beta reputation system. Proceedings of the 15th Bled Electronic Commerce Conference. Bled, Slovenia. 2002. 324-337.
- 3 Hosseini SB, Shojaei A, Agheli N. A new method for evaluating cloud computing user behavior trust. Proceedings of the 7th Conference on Information and Knowledge Technology. Urmia, Iran. 2015. 1-6.
- 4 Liu ZY, Peng DC. User behavior identification for trust management in pervasive computing systems. Proceedings of the 11th IEEE International Workshop on Future Trends of Distributed Computing Systems. Sedona, AZ, USA. 2007. 65-72.
- 5 Noorian Z, Mohkami M, Liu Y, *et al.* SocialTrust: Adaptive trust oriented incentive mechanism for social commerce. Proceedings of 2014 IEEE/WIC/ACM International Joint Conferences on Web Intelligence and Intelligent Agent Technologies. Warsaw, Poland. 2014. 250-257.
- 6 Wang Y, Vassileva J. Bayesian network trust model in peer-to-peer networks. Proceedings of the Second International Workshop on Agents and Peer-to-Peer Computing. Berlin, Heidelberg, Germany. 2004. 23-34.
- 7 Yang RL, Yu XJ. Research on way of evaluating cloud end user behavior's credibility based on the methodology of multilevel fuzzy comprehensive evaluation. Proceedings of the 6th International Conference on Software and Computer Applications. Bangkok, Thailand. 2017. 165-170.
- 8 曹洁, 曾国荪, 姜火文, 等. 云环境下服务信任感知的可信动态级调度方法. 通信学报, 2014, 35(11): 39-49.
- 9 王海艳, 杨文彬, 王随昌, 等. 基于可信联盟的服务推荐方法. 计算机学报, 2014, 37(2): 301-311.
- 10 盛国军, 温涛, 郭权, 等. 基于改进蚁群算法的可信服务发现. 通信学报, 2013, 34(10): 37-48. [doi: 10.3969/j.issn.1000-436x.2013.10.005]
- 11 Liu YZ, Zhang L, Luo P, *et al.* Research of trustworthy software system in the network. Proceedings of the Fifth International Symposium on Parallel Architectures, Algorithms and Programming. Taipei, China. 2012. 287-294.
- 12 Liu ZY, Yau SS, Peng DC, *et al.* A flexible trust model for distributed service infrastructures. Proceedings of the 11th IEEE International Symposium on Object and Component-Oriented Real-Time Distributed Computing. Orlando, FL, USA. 2008. 108-115.
- 13 刘迎春, 郑小林, 陈德人. 基于信任和推荐关系的可信服务发现. 系统工程理论与实践, 2012, 32(12): 2789-2795. [doi: 10.3969/j.issn.1000-6788.2012.12.027]
- 14 Zhang L, Wu J. Research on trustworthy software composition architecture. Proceedings of the 2nd International Conference on Consumer Electronics, Communications and Networks. Yichang, China. 2012. 1273-1276.
- 15 Noorian Z, Marsh S, Fleming M. Prob-Cog: An adaptive filtering model for trust evaluation. Proceedings of the 5th IFIP WG 11.11 International Conference on Trust Management V. Berlin Heidelberg, Germany. 2011. 206-222.
- 16 李小勇, 桂小林. 可信网络中基于多维决策属性的信任量化模型. 计算机学报, 2009, 32(3): 405-416.
- 17 Tian LQ, Ni Y, Lin C. Node behavior trust evaluation based on behavior evidence in WSNs. Proceedings of the 2nd International Conference on Future Computer and Communication. Wuhan, China. 2010. V1-312-V1-317.