

考虑虚拟机间性能互扰基于排队网的多层 Web 应用性能分析模型

杨 雷¹ 代 钰² 张 斌¹ 王 昊¹

(东北大学信息科学与工程学院 沈阳 110819)¹ (东北大学软件学院 沈阳 110819)²

摘 要 多层 Web 应用性能分析是实现资源动态分配和管理、保证多层 Web 应用性能的重要因素之一。传统的多层 Web 应用性能分析模型往往假设服务器部署在无性能互扰的服务器环境中且忽略了逻辑资源服务能力对多层 Web 应用性能的影响。随着云计算的发展,底层物理资源可以通过虚拟化方式形成虚拟资源并对外提供服务,这为多层 Web 应用的性能保证提供了有效支撑。因此,如何考虑虚拟机性能互扰以及逻辑资源服务能力对多层 Web 应用性能的影响已经成为云计算环境中多层 Web 应用性能分析所需解决的关键问题。为此,构建了一个基于排队网的多层 Web 应用性能分析模型,该模型通过丢弃队列来对目前多层 Web 应用性能分析模型在并发数限制方面进行扩展,在考虑虚拟机间性能互扰的基础上,提出了多层 Web 应用性能分析模型参数求解方法。实验结果验证了所提出的多层 Web 应用性能分析模型的有效性。

关键词 性能分析,虚拟机间性能互扰,多层应用

中图分类号 TP393.1

文献标识码 A

DOI 10.11896/j.issn.1002-137X.2015.1.010

Queueing Network Based Performance Model for Multi-tiered Web Application with Consideration of Performance Interference among Virtual Machines

YANG Lei¹ DAI Yu² ZHANG Bin¹ WANG Hao¹

(College of Information Science and Technology, Northeastern University, Shenyang 110819, China)¹

(College of Software, Northeastern University, Shenyang 110819, China)²

Abstract Performance analysis for multi-tiered application is one of the important factors to achieve dynamic resource allocation and management, and insure the better performance of the application. Traditional performance models assume the availability of dedicated hardware for the application and neglect the effect caused by logic resource. With the development of cloud computing, hardware resources can be virtualized as virtual resource to provide service which can support the performance insurance of the multi-tiered application effectively. Then, how to take the performance interference between virtual machines and the logic resource into consideration becomes a key issue when building the performance model for a multi-tiered application. For this, this paper built a performance model for multi-tiered application based on queueing network. This model extends the current performance model with a drop queue to support concurrent connection limit. With the consideration of the performance interference between virtual machines, this paper presented the parameter solving method of the proposed performance model. The experiments show the better performance of the performance model.

Keywords Performance analysis, Performance interference among virtual machine, Multi-tiered application

1 引言

Web 应用已经成为目前的主流网络应用形式之一,多层 Web 应用架构已经广泛应用于电子商务、网上银行等关键业务系统中,其性能保证问题已经成为影响这些业务系统获取商业价值的关键因素之一。由于多层 Web 应用性能受应用负载变化以及资源服务能力的影响^[1],为了保证多层 Web 应

用性能,需要对影响多层 Web 应用性能的因素进行分析,建立一个多层 Web 应用性能与其资源服务能力和负载间的函数关系,以反映应用负载以及资源服务能力变化时多层 Web 应用性能的变化规律。从而,根据该函数关系,通过动态资源分配等手段,对多层 Web 应用的资源服务能力进行动态调整,以适应应用负载的动态变化,进而保证其性能。

然而,目前的研究工作^[2,3]大多未考虑虚拟机间性能互

到稿日期:2013-12-26 返修日期:2014-03-15 本文受国家自然科学基金(60903218,61073062),中央高校基本科研业务费专项资金(N110404008,N110417003,N110804002)资助。

杨 雷(1974—),男,博士,副教授,主要研究领域为性能管理、分布式计算;代 钰(1980—),女,博士,副教授,主要研究领域为性能管理、分布式计算,E-mail: daiyu@mail.neu.edu.cn(通信作者);张 斌(1964—),男,教授,博士生导师,主要研究领域为性能管理、分布式计算;王 昊(1990—),男,硕士生,主要研究领域为性能管理、布式计算。

扰对多层 Web 应用性能的影响,同时在逻辑资源限制情况下的多层 Web 应用性能分析中存在请求丢弃率估计不准确的问题,这将影响多层 Web 应用性能分析的准确性和有效性。为此,本文在目前基于排队论的多层 Web 应用性能分析模型的基础上,对现有模型进行扩展,并且在模型的参数求解过程中考虑了虚拟机间性能互扰对多层 Web 应用性能的影响。

针对于目前多层 Web 应用性能分析所面临的上述问题,本文基于排队网建立了一个多层 Web 应用性能分析模型,该模型通过丢弃队列,对目前多层 Web 应用性能分析模型在并发数限制方面进行扩展,以提高所构建的多层 Web 应用性能分析模型的适用范围,保证多层 Web 应用性能分析的客观性和准确性。通过与目前多层 Web 应用性能分析模型进行对比,验证了本文所提出的多层 Web 应用性能分析模型的有效性。

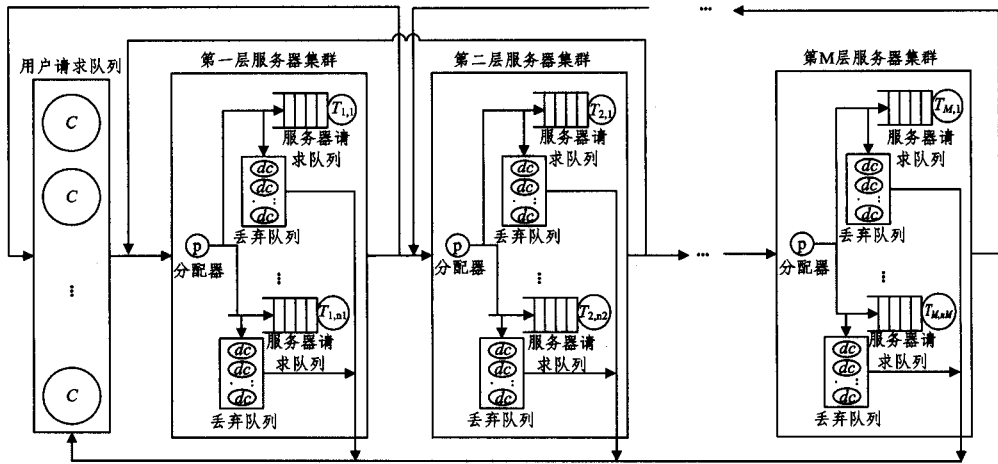


图 1 多层 Web 应用的排队网模型

接下来,本文将根据排队论和效用法则^[4],给出多层 Web 应用端到端响应时间的计算公式。由于 CPU 资源消耗与请求到达率之间具有线性关系,而 I/O 类资源的消耗与请求到达率之间是非线性关系,因此在下面的讨论中将请求在每层的平均停留时间计算分成 2 个部分:请求在每层 CPU 资源上的平均停留时间以及请求在每层非 CPU 资源上的平均停留时间。在下文中,本文将先讨论一般队列系统的响应时间计算公式,并在此基础上给出多层 Web 应用的端到端响应时间计算公式。

$$R_{cpu}^{i,j} = \sum_k \frac{1}{\lambda_j \times (1 - p_{j,k}^{drop})} \times \frac{f_{cm}(\lambda_{1,j} \times (1 - p_{1,k}^{drop}) \times p_j^k, \lambda_{2,j} \times (1 - p_{2,k}^{drop}) \times p_j^k, \dots, \lambda_{i,j} \times (1 - p_{i,k}^{drop}) \times p_j^k, CPU_j^k)}{T} \quad (1)$$

这里, $u_{j,k}^{cpu}$ 为在时间间隔 T 内层 j 的服务器 k 的 CPU 利用率; $p_{j,k}^{drop}$ 为层 j 服务器 k 的平均请求丢弃概率; p_j^k 为发送到层 j 的请求被分配给服务器 k 进行处理的概率; $\lambda_{u,j}$ 为层 j 事务类型为 u 的请求到达率。

这样,可以得到请求在所有层的平均停留时间之和,即请求在多层 Web 应用的端到端响应时间。式(2)给出了一个多层 Web 应用端到端响应时间计算公式。

$$R = \sum_{j=1}^M R_{cpu}^{i,j} + \sum_{j=1}^M \sum_{k=1}^{n_j} O_{j,k}^i \quad (2)$$

这里, $O_{j,k}^i$ 是在层 j 的虚拟服务器 k 中在请求负载向量为 λ_j 时由虚拟机间性能互扰所导致的请求平均延迟时间。

在某一时间间隔 T 内多层 Web 应用 App 的应用负载向量 $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_n]$ 已知的情况下,要计算得到端到端响应

2 基于排队网的多层 Web 应用性能建模以及参数求解方法

排队网是由一系列排队节点构成的网络,每个排队节点是一个服务站。本文采用 M/G/PS 型排队网络模型分析多层 Web 应用的性能,这里每个服务器将以服务器共享方式(processor sharing)调度并处理用户请求。

在一个多层 Web 应用中,每一层为其上一层提供功能服务,并调用其下一层功能,以完成请求处理过程。以下本文将先分析多层 Web 应用的请求处理过程,并在此基础上,基于排队论对多层 Web 应用性能进行建模。一个多层 Web 应用可以建模为一个如图 1 所示的排队网模型,其中: C 表示客户请求; P 表示分配器; T_{ij} 表示第 i 层的服务器 j ; dc 表示丢弃的用户请求。

在时间段 T , CPU 利用率可以近似地看作是 CPU 忙时间与 T 的比值,即 CPU 消耗与 T 的比值。这里, CPU 消耗是指处理请求所占用的 CPU 服务时间之和。在系统中有不同应用负载的情况下,即存在应用负载密度向量 $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_i]$ 的情况下,不同事务类型请求具有不同的 CPU 消耗,即其 CPU 服务时间不同, CPU 消耗与负载密度向量 $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_i]$ 间存在某一函数关系。记该函数关系为 $CM = f_{cm}(\lambda = [\lambda_1, \lambda_2, \dots, \lambda_i, CPU_j^k])$, 由此对于层 j , 请求在该层 CPU 资源上的平均停留时间可由式(1)计算。

时间 R , 需要计算以下参数的取值:

- 到达某一层 j 的应用负载向量。 $\lambda_j = [\lambda_{1,j}, \lambda_{2,j}, \dots, \lambda_{n,j}]$, 本文假设: 层与层之间是线性关系, 到达某一层 j 的应用负载向量应该为其上一层未丢弃应用负载向量, 即 $\lambda_j = \lambda_{j-1}$, 且 $\lambda_1 = \lambda_0 = \lambda$ 。

- 层 j 的服务器 k 的请求丢弃率 $p_{j,k}^{drop}$ 。某一时刻 t , 层 j 的服务器 k 的请求丢弃概率与该服务器在当前时刻所到达的应用负载(包括应用负载大小以及类型)有关。为此, 本文使用状态空间模型来计算请求丢弃率。

- 到达层 j 的请求分配给服务器 k 进行处理的概率 p_j^k 。将其看作是一个历史经验值。

- 层 j 的服务器 k 的 CPU 资源的应用负载向量为 $\lambda_j = [\lambda_{1,j}, \lambda_{2,j}, \dots, \lambda_{n,j}]$ 时的 CPU 消耗。将在下面给出其计算方

法。

- 每层的平均延迟时间 O_{ij}^k 。将在下面给出其计算方法。

法。

可以建立一个如式(3)所示的 CPU 消耗与负载向量间的关系函数。

$$C_m = a_0 + a_1 \lambda_1 + a_2 \lambda_2 + \dots + a_n \lambda_n + \gamma_1 \lambda_1^2 + \gamma_2 \lambda_2^2 + \dots + \gamma_n \lambda_n^2 \quad (3)$$

根据式(3),CPU 消耗与请求负载向量 $\lambda^T = [\lambda_1, \lambda_2, \dots, \lambda_n]^T$ 间存在非线性关系。这样,可以利用非线性回归分析方法对式(3)中的系数进行求解。限于篇幅,本文将不在此进行讨论。

记要建模的虚拟机为 k ,本文将建立当前虚拟机 CPU 利用率 $CPUu_k$ 、背景虚拟机 CPU 利用率 $CPUu_b$ 、当前虚拟机读请求数 ior_k 、当前虚拟机写请求数 iow_k 、背景虚拟机读请求数 ior_b 和背景虚拟机写请求数 iow_b 与响应时间延迟 o 之间的关系,如式(4)所示。

$$o = a_0 + a_1 \times CPUu_k + a_2 \times CPUu_b + a_3 \times ior_k + a_4 \times iow_k + a_5 \times ior_b + a_6 \times iow_b \quad (4)$$

利用线性回归方法对系数 $a_0, a_1, \dots, a_6$ 求解的基本思想是:求解系数 $a_0, a_1, \dots, a_6$,使得在该系数下,估计值 o' 与实际观测值 o 的残差平方和最小。在求得系数 $a_0, a_1, \dots, a_6$ 后,可以利用式(4)对由虚拟机间性能互扰所导致的平均请求延迟时间进行估计。进而,能够根据式(2)对多层 Web 应用端到端响应时间进行计算。

3 实验对比

本节将给出对虚拟服务器环境下的多层 Web 应用端到端平均响应时间进行估计的实验结果。在实验中,Rubis 的 apache 端包含 2 台虚拟服务器,database 端包含 1 台虚拟服务器。每台虚拟服务器的配置相同。

利用本文所提出的考虑虚拟机性能互扰的多层 Web 应用端到端响应时间估计方法对多层 Web 应用响应时间进行计算,比较端到端响应时间估计值与实际值间的差异。实验中,database 端最大并发数限制设置为 30,apache 端最大并发

数限制设置为 100。在请求到达率大小一定时,在请求事务类型混合比例不同的情况下分别进行 20 次实验,取响应时间的平均值。实验结果如图 2 所示。

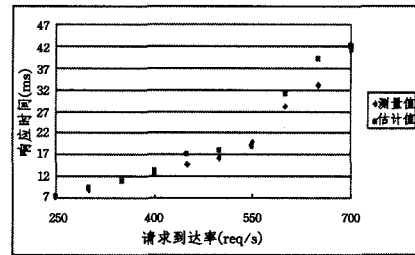


图 2 虚拟服务器环境下多层 Web 应用响应时间估计对比

图 2 说明了在大多数情况下利用本文方法估计得到的响应时间与实际测量值相比误差并不大。

结束语 本文建立了一个基于排队网的多层 Web 应用性能分析模型,给出了多层 Web 应用端到端响应时间计算公式,并给出了请求丢弃率估计算法、不考虑虚拟机间性能互扰的多层 Web 应用端到端响应时间计算方法以及考虑虚拟机间性能互扰的多层 Web 应用端到端响应时间计算方法。实验验证了本文所提出的多层 Web 应用性能分析模型的有效性。

参考文献

- [1] Cook B, Babu S, Candea G, et al. Toward Self-Healing Multitier Services[C]// ICDE Workshops (SMDB 2007). Washington, 2007:424-432
- [2] Zhang Qi, Cherkasova L, Smirni E. A Regression-Based Analytic Model for Dynamic Resource Provisioning of Multi-Tier Applications[C]// ICAC'07. 2007;27-37
- [3] Shoaib Y, Das O. Web application performance modeling using layered queueing networks[J]. Electronic notes in theoretical computer science, 2011,275:123-142
- [4] Lazowska E D, Zahojan J, Graham G S. Quantitative System Performance: Computer System Analysis using Queueing Network Models[M]. Upper Saddle River, Prentice-Hall, 1984
- [5] IEEE. Reference on Service Science and Innovation (ICSSI). IEEE, 2013: 50-53
- [6] Zhang J, Tao Y. The Interaction based Innovation Process of Architectural Design Service[C]// IEEE International Conference on Industrial Engineering and Engineering Management, 2007:1719-1723
- [7] Bailey M W, VerDuin W H. FIPER: an intelligent system for the optimal design of highly engineered products[C]// Measuring Performance and Intelligence of Intelligent Systems, 2001:467-477
- [8] Goel S, Talya S, Sobolewski M. Preliminary Design Using Distributed Service-based Computing[C]// Proceeding of the 12th Conference on Concurrent Engineering: Research and Applications, ISPE, Inc., 2005:113-120
- [9] Chen C H. Constructing intelligent design platform for co-design novel interior design service[C]// 2013 Fifth International Conference on Service Science and Innovation (ICSSI). IEEE, 2013: 50-53
- [10] Cornford S, Cole B, Dubos G, et al. Evaluating fractionated space systems-Status[C]// Aerospace Conference 2013. IEEE, 2013:1-10
- [11] 马继楠, 刘莉, 龙腾. 基于 MODELCENTER 的 EPUAV 系统模型建立与参数优化[J]. 系统仿真学报, 2009,16:5203-5206
- [12] 张晓丽, 迟国政, 朱鸣华, 等. 基于 Web 的工程计算平台研究[J]. 大连理工大学学报, 2004(2):259-261
- [13] 刘向峰, 胡荣, 高志. 基于 Web 环境的机械设计的异地协同设计[J]. 机械设计, 2002(12):23-25
- [14] Sobolewski M, Kolonay R. Unified Mogramming with Var-Oriented Modeling and Exertion-Oriented Prgramming Languages [J]. International Journal of Communications, Network and System Sciences, 2012(5):579-592

(上接第 22 页)