

DOI: 10.7652/xjtuxb201606005

多版本视频点播流媒体服务器集群资源分配方法

赵辉^{1,2}, 郑庆华^{1,2}, 张未展^{1,2}, 李珍艳^{1,2}, 叶舒雁^{1,2}

(1. 西安交通大学计算机科学与技术系, 710049, 西安; 2. 陕西省天地网技术重点实验室, 710049, 西安)

摘要: 针对大规模移动学习应用背景下流媒体服务器集群资源分配过量或不均的问题,提出了一种面向移动学习系统的多版本视频点播流媒体服务器集群资源分配方法。该方法通过用户历史点播行为日志分析,挖掘在不同终端环境下用户点播行为的特征和规律,在此基础上,采用排队论理论进行多版本视频点播中流媒体服务器集群资源分配建模,并通过实时预测用户请求到达率的变化情况,动态调整资源的分配,从而实现了集群资源的动态配置。实验结果表明,所提方法的平均服务拒绝率和资源利用率分别在 1% 和 80% 左右,既保证了用户体验满意度,又降低了系统服务成本。

关键词: 移动学习系统;多版本视频点播;用户行为分析;排队论;资源分配

中图分类号: TP391 **文献标志码:** A **文章编号:** 0253-987X(2016)06-0030-06

A Resource Allocation Method for Clusters of Streaming Media Servers in Multi-Version VoD

ZHAO Hui^{1,2}, ZHENG Qinghua^{1,2}, ZHANG Weizhan^{1,2}, LI Zhenyan^{1,2}, YE Shuyan^{1,2}

(1. Department of Computer Science and Technology, Xi'an Jiaotong University, Xi'an 710049, China;
2. Shaanxi Province Key Laboratory of Satellite and Terrestrial Network Tech. R&D, Xi'an 710049, China)

Abstract: A resource allocation method for clusters of streaming media servers in multi-version VoD is proposed to solve the problem of the inappropriate resources allocation for clusters of streaming media servers in large-scale mobile learning (m-learning) system. The method analyzes user's historical learning logs and mines user's behavior characteristics, and then it utilizes the queueing theory to establish a resource allocation model for the cluster of streaming media servers. The method uses the model to predict the user arrival rate in real-time and then allocates appropriate resources dynamically. Simulation results show the correctness and efficiency of the proposed method. The average service rejection rate and average resource utilization rate are around 1% and 80% respectively. The proposed method can ensure the user experience satisfaction and reduce service cost.

Keywords: m-learning; multi-version VoD; user behavior analysis; queueing theory; resource allocation

随着移动通信技术以及智能终端的发展,人们 可以在移动设备上访问视频资源,推动了传统的数

字化学习(e-Learning)向移动学习(m-Learning)的转变。由于终端分辨率、屏幕大小的不同和网络的不同,需要为用户提供不同码率、分辨率的视频流,即多版本视频点播。多版本视频点播的研究有以下几种:①可分级视频编码技术^[1];②实时视频转码^[2];③预存储多版本视频^[3];④平衡视频存储和转码^[4],少量的热点视频存储多个版本,大量的非热点视频存储高质量版本,通过实时转码技术提供低质量版本的点播。

正是在这种平衡视频存储与转码的多版本视频点播背景下,移动学习系统需要分配转码计算资源和带宽资源。本文提出了一种面向移动学习系统的多版本视频点播流媒体服务器集群资源分配方法。首先,分析移动学习用户历史点播行为日志,挖掘用户点播行为的特征和规律,建立用户视频点播行为模型。其次,采用排队论理论进行多版本视频点播流媒体服务器集群资源分配建模,并根据用户到达率的变化情况动态地调整集群资源的分配。

1 相关工作

目前资源分配方法的研究分为:基于控制论的方法、基于性能模型的方法及基于负载预测的方法。

(1)基于控制论的方法。文献[5]开发了一个管理框架为虚拟化应用自动分配资源。文献[6]通过经典反馈控制论为系统提供服务质量保证的资源分配方法。它们都依靠底层服务器系统的性能,适用于面向网页链接等短时占用资源的应用,不适用于流媒体点播持续性占用资源的应用。

(2)基于性能模型的方法。文献[7]提出基于群体智能的方法来决定服务器资源分配策略和作业调度方法,将资源分配问题模拟为生态现象,采用将混合离散粒子群体优化与进化博弈论相结合的方法选择性能适合的虚拟机。文献[8]提出了一种动态资源缩放方案,利用前端负载平衡器统计活动会话,根据当前在线用户活动会话的数量分配资源,在资源自动配置的基础上根据当前活动会话的数量动态缩放虚拟机资源。由于流媒体点播需求具有较大的波动性,这类方法难以确定流媒体应用中实时的性能阈值,无法准确地分配所需的资源。

(3)基于工作负载预测的方法。文献[9]预测每个应用程序将来的负载情况,以及预测每一个虚拟机将来的性能,然后分配每个应用的资源量。文献[10]通过跟踪分析数据中心的应用运行日志,提出基于马尔可夫模型的工作负载模式,以时间相关性

为特点,描述了虚拟机集群与工作负载模式的变化规律,进而根据工作负载模型预测应用所需资源。这类方法的主要挑战在于如何确定良好的预测模型,提出的工作负载模型复杂度特别高,不容易应对用户请求的实时变化,无法体现移动学习视频点播应用类型的工作负载规律。

综上,已有的资源分配方法研究均未考虑移动学习系统及多版本视频点播系统的特殊性,也未考虑如何利用用户行为分析来分配资源。

2 流媒体服务器集群资源分配方法

视频点播服务过程如图 1 所示,一般可描述为排队系统^[11]。本文通过对历史点播行为日志进行分析,得到排队论所需的各个参数。

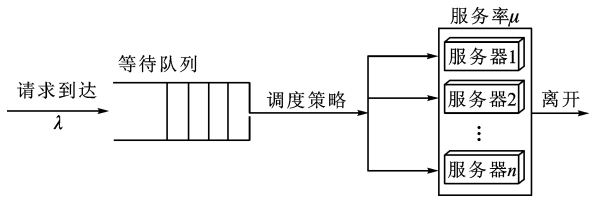


图 1 视频点播服务过程图

2.1 用户点播行为分析

本文采用西安交通大学网络教育学院移动学习系统的日志数据,选取了从 2014 年 12 月 1 日至 2015 年 1 月 31 日期间产生的所有日志文件,共计 0.98 GB。数据预处理后共有日志 92 632 条,其中电脑端日志 66 424 条,移动端日志 26 208 条,总课程数 5 142 节。

用户请求到达率:选取若干天的点播日志进行统计,将 1 min 作为单位时间,计算单位时间内用户请求到达个数,统计到达个数出现的概率。图 2 为统计的用户请求到达率分布,可以看出用户请求到达率的分布近似于泊松分布^[12] $P(X=i)=(\lambda^i/i!)\cdot e^{-\lambda}, i=0,1,2,\dots$ 。

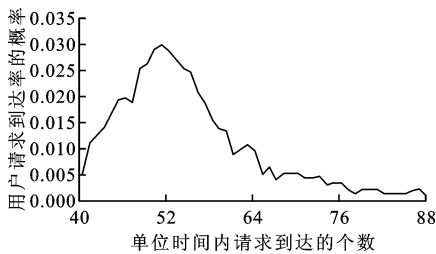


图 2 用户请求到达率的分布图

视频点播热度分布为通过统计点播日志得出的所有视频的点播热度分布,呈现出明显的“长尾”效

应,如图 3 所示。相关研究多使用 Zipf-like 分布^[13]描述为

$$p_i = s_i / S = (i^\theta \sum_{i=1}^M i^{-\theta})^{-1}$$

式中: p_i 表示视频 i 的点播概率; s_i 为视频 i 的点播次数; S 为所有视频的点播总次数; M 为视频总数; θ 为分布参数。

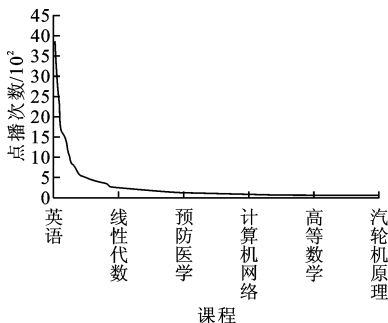


图 3 视频点播热度分布图

不同版本间的点播热度分布:在多版本视频点播的应用背景下,还需要考虑同一个视频的不同版本之间的点播概率。具体概率可通过统计同一个视频在不同终端的点播次数计算得到,文献[14]使用高斯分布表示为 $p'_j = e^{-(j-m)^2/2\sigma^2} / ((2\pi)^{1/2} \sigma)$ 。其中: j 为视频版本号; p'_j 为版本 j 的点播概率; m 为高斯分布位置参数; σ 为高斯分布尺度参数, σ 越小,表明在 m 处的概率越大。

服务时长分布为针对不同课程视频分别统计不同终端的学习日志,得到学习时长。部分课程的学习时长统计如图 4 所示。由于各个用户请求所需的服务时长相互独立,不同课程视频的学习时长可用一般分布 $G(t), t \geq 0$ 描述,即用户的服务时长独立,平均服务时长为 $0 < t = 1/\mu = \int_0^\infty t dG(t)$ 。

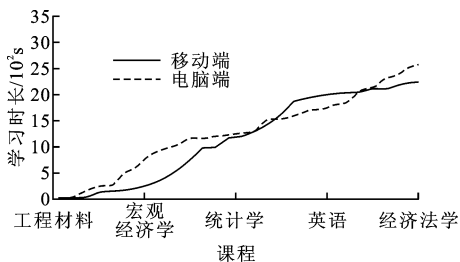


图 4 学习时长的分布

单个请求所占带宽资源可根据每个视频的每个版本的点播概率以及相应的视频码率计算得到,一个点播请求所需占用的平均带宽资源可表示为 $R = E[r_i] = \sum p_{i,j} r_{i,j}$, 其中 $p_{i,j}$ 为第 i 个视频的第

j 个版本的点播概率, $r_{i,j}$ 为其码率。

对于实时转码,不同码率、分辨率的视频之间转码所需的计算资源不同,版本 i 转到版本 j 的转码权重为 $w_{i,j}$, 可通过统计 CPU 利用率得到。所需的转码计算资源 $U = E[w_i] = \sum p_{i,j} w_{i,j}$, 则单位 CPU 能够同时并发执行的转码任务数为 $V=1/U$ 。

用户平均等待时长指从用户发出点播请求开始,到服务器分配资源开始服务为止的这段时长。本文中,用户平均等待时长是已知约束条件。

2.2 集群资源分配方法

采用排队论模型建立集群资源分配的思路为:首先,训练不同的用户请求到达率情况下的资源需求量;其次,预测用户请求到达率的变化情况;最后,根据训练结果和用户请求到达率,实时动态调整资源的分配。

2.2.1 排队论模型训练 使用排队论将该问题描述为用户请求到达率服从参数为 λ 的泊松分布,实时转码的请求数占所有到达的请求数的比例为 ρ (转码部分请求到达率服从参数为 $\rho\lambda$ 的泊松分布),集群服务率为 μ , 请求服务时间是平均值为 $1/\mu$ 的一般分布,一个点播请求所需占用的平均带宽资源为 R , 单位 CPU 能够同时开启的转码任务数为 V , 用户平均等待时长为 T 时,计算至少需要分配的带宽资源和计算资源。可用 $M/G/n$ 排队论模型^[15]描述该问题,则请求等待概率为

$$P_Q = \frac{(\rho\lambda)^n}{n!(1-\rho)} \left[\sum_{k=0}^{n-1} \frac{(\rho\lambda)^k}{k!} + \frac{(\rho\lambda)^n}{n!(1-\rho)} \right]^{-1} \quad (1)$$

队列中正在等待的请求平均个数(队列长度)

$$N_Q = \rho P_Q \bar{X}^2 / (2\bar{X}^2(1-\rho)) \quad (2)$$

式中: $\rho = \lambda/n\mu$, n 为服务器并行服务能力; $\bar{X}^2 = [t_{ij}^2] = \sum p_{ij} t_{ij}^2$ 为请求服务时长方差。队列中的平均等待时长为 $T = N_Q/\lambda$ 。

给定 $M/G/n$ 排队论模型相关参数后,分别采用逐步逼近法求出所需的服务器并行流化服务能力 n_1 和并行转码能力 n_2 。一个点播请求所需占用的平均带宽资源为 R , 单位 CPU 能够同时开启的转码任务数为 V , 故所需的带宽资源为 $B = n_1 R$, 所需的转码计算资源为 $C = n_2 / V$ 。

2.2.2 用户请求到达率预测 本文采用二次指数平滑法来预测用户请求到达率。用 $F_t^{(1)}$ 和 $F_t^{(2)}$ 表示 t 时刻的一次指数平滑值和二次指数平滑值, $F_t^{(1)} = \alpha Y_t + (1-\alpha) F_{t-1}^{(1)}$, $F_t^{(2)} = \alpha F_t^{(1)} + (1-\alpha) F_{t-1}^{(2)}$, 其中, Y_t 是 t 时刻的用户请求到达率, α 是平滑系数

($0 < \alpha < 1$)。 t 时刻的一次、二次指数平滑值计算出来后,则 $t + \Delta t$ 时刻的预测值为 $F_{t+\Delta t} = a_t + b_t \Delta t$, 其中 $a_t = 2F_t^{(1)} - F_t^{(2)}$, $b_t = \alpha(F_t^{(1)} - F_t^{(2)}) / (1 - \alpha)$ 。

2.2.3 动态资源分配 若用户点播请求持续增加,当用户点播请求到达率达到某种情况时,需及时地新增虚拟机,使其加入该集群提供服务。若用户点播请求持续降低,集群资源利用率也会随之降低,则需要回收虚拟机,其思路如下:

步骤 1 维护最近 n 个统计时间段内的请求到达率;
步骤 2 若当前用户请求到达率达到系统所能服务的最大到达率的 $1/2$ 时,或当前用户请求到达率降低到集群减少一台虚拟机之后仍能满足服务时,根据当前保存的请求到达率集合,利用二次指数平滑预测法来预测用户请求到达率;
步骤 3 若请求到达率呈降低趋势,则需回收一台虚拟机,使其撤出服务集群,此时选择资源利用率最低的虚拟机作为待回收虚拟机提交给调度节点,调度节点将该虚拟机已有的请求迁移至其他虚拟机后,回收该虚拟机,并将其移出集群;
步骤 4 若请求到达率呈现增长趋势,则预测一段时间(虚拟机启动延迟)后的用户请求到达率,若其不小于集群所能承载的最大到达率,则新增虚拟机,并将其加入流媒体服务器集群。

3 实验结果与分析

实验数据采用西安交通大学网络教育学院移动学习系统的日志数据,实验中的参数都是通过用户点播历史日志分析所得。本文从资源利用率和资源拒绝率两个指标验证分配的资源是否合适。

带宽资源利用率 $U_b = b/B$, 计算资源利用率 $U_c = c/C$ 。其中: b 为已消耗的带宽, B 为分配的总带宽; c 为执行转码任务的 CPU 核数, C 为分配的 CPU 总核数。

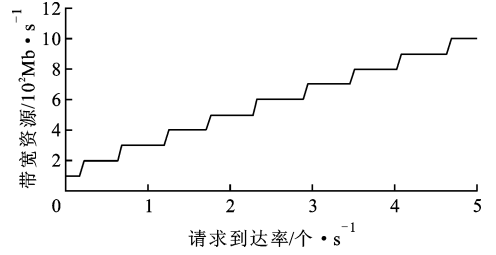
服务拒绝率 J 为视频点播过程中由于请求在队列中的等待时长超过最大限定值 T_{wait} 而被系统拒绝的请求数量与到达的总请求数量之比,计算公式为 $J = n_{\text{reject}} / n_{\text{total}}$ 。

3.1 排队论模型训练结果分析

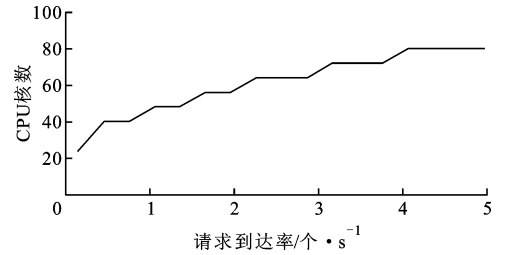
在第一个实验中,首先限定请求等待时长为 5 s,使用排队论模型分别训练在不同的用户请求到达率(个/s)的情况下,需要为系统分配的带宽资源和 CPU 计算资源,结果如图 5 所示。在第 2 个实验中,设置用户请求到达率为 1 个/s,通过改变平均等待时长从 1 s 到 59 s,分别训练得到为系统分配的

带宽资源和 CPU 计算资源,结果如图 6 所示。

由图 5、图 6 可以看出:随着请求到达率的增高,所需的带宽资源和计算资源都逐步增大;随着平均等待时长的增大,所需的资源逐步减少。

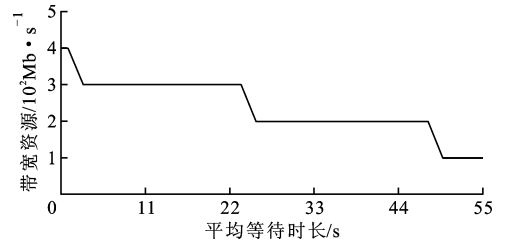


(a) 请求到达率与带宽资源的关系

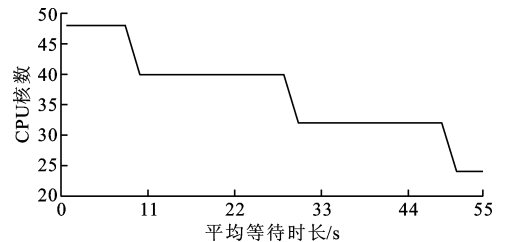


(b) 请求到达率与 CPU 资源的关系

图 5 请求到达率与分配资源的关系



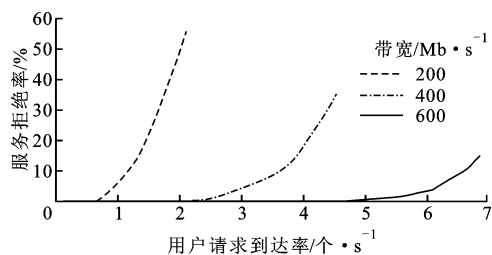
(a) 平均等待时长与带宽资源的关系



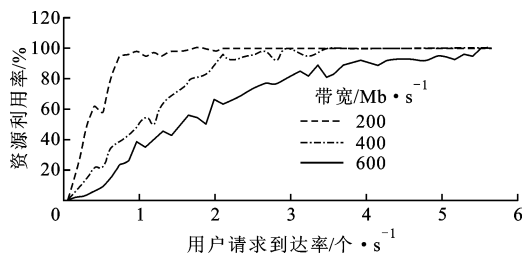
(b) 平均等待时长与 CPU 资源的关系

图 6 平均等待时长与分配资源的关系

本文使用服务拒绝率与资源利用率来验证资源分配是否合理。以带宽资源为例,选取集群带宽为 200、400、600 Mb/s,并限定平均等待时长为 5 s 时,服务拒绝率和资源利用率的情况如图 7 所示。由图 7 可以看出:在系统带宽资源为 200 Mb/s 时,当请求到达率小于 0.5 个/s,系统服务拒绝率较低,当请求到达率超过 0.5 个/s 时,系统服务拒绝率开始出

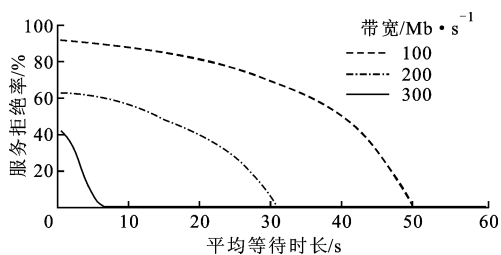


(a) 用户请求到达率与服务拒绝率的关系

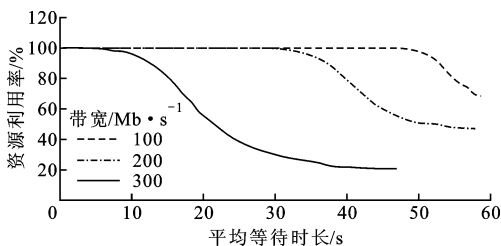


(b) 用户请求到达率与资源利用率的关系

图7 不同用户请求到达率下的服务拒绝率和资源利用率



(a) 平均等待时长与服务拒绝率的关系



(b) 平均等待时长与资源利用率的关系

图8 不同平均等待时长的服务拒绝率和资源利用率

现,并呈现快速上升趋势;同样,当用户请求到达率小于 0.5 个/s 时,资源利用率未到 100%,而当请求到达率超过 0.5 个/s 时,资源利用率几乎达到 100%。如图 5a 所示,带宽为 200 Mb/s 可服务的请求到达率最大值为 0.62,与图 7 中服务拒绝率和资源利用率出现拐点所对应的请求到达率基本一致。同样,当带宽为 400、600 Mb/s 时,服务拒绝率和资源利用率出现拐点所对应的请求到达率与图 5a 中的结果也基本一致,说明本文分配的资源是合适的。

若限定请求到达率为 1 个/s,请求队列平均等待时长 t 从 1 s 到 60 s,当带宽分别为 100、200、300 Mb/s 时,服务拒绝率和资源利用率情况如图 8 所示。可以看出,当带宽为 100、200 和 300 Mb/s 时,为使系统的服务拒绝率为 0,则需要最小的平均等待时长分别为 50、32 和 6 s,而图 6a 中 100、200 和 300 Mb/s 带宽可满足的平均等待时长分别为 52、35 和 5 s,它们基本上保持一致,说明我们分配的资源是合适的。

3.2 动态资源分配结果分析

用户的点播次数随时间具有较大的波动性,随着用户请求到达率的变化,系统需要进行动态资源分配,以实时调整集群资源量。图 9 为某 50 h 随着请求到达率的变化而动态调整集群资源量的变化情况,图 10 为对应的服务拒绝率和资源利用率。

由图 10 可见,随着用户请求到达率的变化,系统的服务拒绝率基本都很低,同时资源利用率基本在 60% 以上,平均拒绝率和资源利用率在 1% 和

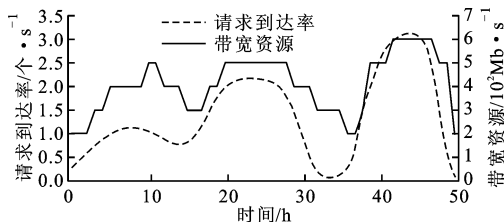


图9 动态带宽分配

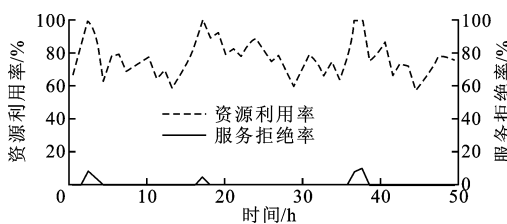


图10 动态资源利用率和服务拒绝率

80% 左右。当请求到达率经过 3 个持续增大与减小过程时,系统带宽资源量呈现出相同的变化趋势。资源利用率在 3 个时刻出现 100% 的情况,而在相同的时刻,也出现了服务拒绝率 10% 左右的现象,这是因为在这 3 个时刻请求到达率持续变大,而此时新分配的资源还未加入服务。同理,在请求到达率持续变小时,可以看到资源利用率和服务拒绝率都比较低,这是由于分配的过多的资源还未及时回收。

4 结 论

本文以移动学习为应用背景,针对多版本视频点播系统提出了一种面向移动学习系统的多版本视频点播流媒体服务器集群资源分配方法。通过分析用户点播日志数据,挖掘用户点播行为规律与特征,

建立用户点播行为模型,采用排队论模型来训练不同的用户请求到达率情况下所需的资源,并实时预测用户请求到达率变化情况,从而动态地为集群分配资源。最后,通过资源利用率和服务拒绝率两个指标来验证了所提方法的正确性和有效性。下一步我们将进一步挖掘移动用户的行为特征,并丰富我们的资源分配方法。

参考文献:

- [1] ZHU Zuqing, LI Suoheng, CHEN Xiaoliang. Design QoS-aware multi-path provisioning strategies for efficient cloud-assisted SVC video streaming to heterogeneous client [J]. *IEEE Transactions on Multimedia*, 2013, 15(4): 758-768.
- [2] HSU T, LI Y. A weighted segment-based caching algorithm for video streaming objects over heterogeneous networking environments [J]. *Expert Systems with Applications*, 2011, 38(4): 3467-3476.
- [3] MA Xiaoqiang, WANG Haiyang, LI Haitao, et al. Exploring sharing patterns for video recommendation on YouTube-like social media [J]. *Multimedia Systems*, 2014, 20(6): 675-691.
- [4] ZHAO Hui, ZHENG Qinghua, ZHANG Weizhan, et al. A version-aware computation and storage trade-off strategy for multi-version VoD systems in the cloud [C]// *Proceedings of the 20th IEEE Symposium on Computers and Communication*. Piscataway, NJ, USA: IEEE, 2015: 943-948.
- [5] DUTREILH X, RIVIERRE N, MOREAU A, et al. From data center resource allocation to control theory and back [C]// *Proceedings of the 3rd IEEE International Conference on Cloud Computing*. Piscataway, NJ, USA: IEEE, 2010: 410-417.
- [6] PAN Wenping, MU Dejun, WU Hangxing, et al. Feedback control-based QoS guarantees in web application servers [C]// *Proceedings of the 10th IEEE International Conference on High Performance Computing and Communications*. Piscataway, NJ, USA: IEEE, 2008: 328-334.
- [7] LEMBOUCHER C, CHELOUAH R, SIARRY P, et al. A swarm intelligence method combined to evolutionary game theory applied to the resources allocation problem [J]. *International Journal of Swarm Intelligence Research*, 2012, 3(2): 20-38.
- [8] HUBER N, BROSIG F, KOUNEV S. Model-based self-adaptive resource allocation in virtualized environments [C]// *Proceedings of the 6th International Symposium on Software Engineering for Adaptive and Self-Managing Systems*. Piscataway, NJ, USA: IEEE, 2011: 90-99.
- [9] ARDAGNA D, GHEZZI C, PANICUCCI B, et al. Service provisioning on the cloud: distributed algorithms for joint capacity allocation and admission control [C]// *Proceedings of the 3rd European Conference on Towards a Service-Based Internet, Service Wave 2010*. Berlin, Germany: Springer, 2010: 1-12.
- [10] KHAN A, YAN Xifeng, TAO Shu, et al. Workload characterization and prediction in the cloud: a multiple time series approach [C]// *Proceedings of the IEEE Network Operations and Management Symposium*. Piscataway, NJ, USA: IEEE, 2012: 1287-1294.
- [11] NAN Xiaoming, HE Yifeng, GUAN Ling. Queueing model based resource optimization for multimedia cloud [J]. *Journal of Visual Communication and Image Representation*, 2014, 25(5): 928-942.
- [12] 凌强, 张逸成, 严金丰, 等. 视频点播系统用户行为模型的构建与应用 [J]. *小型微型计算机系统*, 2013, 34(3): 548-552.
- LING Qiang, ZHANG Yicheng, YAN Jinfeng, et al. Construction and application of users' behavior model in the video on demand system [J]. *Journal of Chinese Computer Systems*, 2013, 34(3): 548-552.
- [13] CIULLO D, MARTINA V, GARETTO M, et al. How much can large-scale Video-on-Demand benefit from users' cooperation? [C]// *Proceedings of the IEEE INFOCOM*. Piscataway, NJ, USA: IEEE, 2013: 2724-2732.
- [14] SHEN B, LEE S J, BASU S. Caching strategies in transcoding-enabled proxy systems for streaming media distribution networks [J]. *IEEE Transactions on Multimedia*, 2004, 6(2): 375-386.
- [15] 曹永荣, 胡伟. 基于现场服务排队近似 M/G/m 模型的 CSR 配置 [J]. *重庆师范大学学报: 自然科学版*, 2010, 4: 36-40.
- CAO Yongrong, HU Wei. Customer service representative staffing based on after-sales field service queuing approximation M/G/m model [J]. *Journal of Chongqing Normal University: Natural Science*, 2010, 4: 36-40.

(编辑 武红江)