

通用链路层上实时业务的多接入分组调度^①

崔 扬^② 徐玉滨 许荣庆 沙学军 丁 哲

(哈尔滨工业大学通信技术研究所 哈尔滨 150080)

摘 要 提出了一个能够为在通用链路层(GLL)上的实时分组业务提供 QoS 保障的多接入分组调度算法。首先设计了一个目标是在保障实时业务 QoS 的同时避免资源浪费的调度效用函数,该函数表示将一个用户数据包调度到一个无线接入链路上进行传输所带来的效用;其次在这个调度函数的基础之上兼顾了公平性,建立一个多接入分组调度模型,值得注意的是该模型是一个 NP 问题;最后利用 Hopfield 神经网络来快速有效地为这个调度模型找出优化解。仿真结果表明,与 M-LWDF 和 PLR 等典型算法相比,该算法在高系统负载的情况下能够满足实时分组的时延要求,同时提高了频谱效率并降低了丢包率及时延抖动。

关键词 通用链路层(GLL),多接入分组调度,调度效用函数,神经网络

0 引 言

由于无线通信技术的快速发展及受产业利益的影响,目前无线接入领域中出现了大量使用不同无线接入技术和不同覆盖面积的无线接入网络,例如使用 WCDMA、TD-SCDMA 和 CDMA2000 的蜂窝网以及无线宽带通信网中的 WiMAX、WLAN 及 WiBro^[1]。因此如何融合这些异构的无线接入网络来实现无线资源的有效利用将成为一种不可避免的发展趋势,这就给网络设计者在怎样协同地管理这些异构接入网络上带来了巨大的挑战。在这一背景下,通用链路层(general link layer, GLL)的引入将会为应对这一挑战提供有效的方法^[2]。GLL 位于网络中的层 2(接入链路层)与层 3(网络层)之间,并在属于不同无线接入技术的无线链路之间进行实时、统一的链路处理。它的一个主要功能就是根据信道质量标识将来自上层的数据流有效分配到各个独立的无线接入链路(radio access, RA)上,并完成相应的协议转化。许多文献已经证明了多接入能够带来较大的系统收益,如文献[3]使用贪婪算法获得了最高 15% 的频谱增益。同时应注意到,未来的网络将是基于分组交换的,电路交换将会最终退出。

例如在 WiMAX^[4] 和 3GPP 的长期演进^[5] 的网络都是基于 IP 的。因此未来的异构网络必然是由基于 IP 的分组交换网络所构成,这样一来运行其上的实时分组业务,如 VoIP,必将成为主要业务。因此怎样快速有效地在不同的 RA 之间进行分组调度来为实时业务提供 QoS 保障将至关重要。由于无线信道的时变特性及频谱资源的有限性,算法的执行时间以及频谱效率都应当在分组调度算法中考虑到。目前很多文献已经提出了针对无线网络中实时业务 QoS 保障的分组调度算法,例如 M-LWDF 算法^[6] 和 PLR 算法^[7]。然而这些算法均没有考虑到频谱效率和调度的执行时间。针对上述问题,本文提出了一个多接入分组调度算法 QU 来为 GLL 上的实时分组业务提供 QoS 保障。与 M-LWDF 和 PLR 算法相比,在负载较高的情况下,本文所提算法能够在满足实时业务时延需求的同时提高频谱效率、降低丢包率并减少时延抖动。

1 系统描述

图 1 给出了实时业务多接入调度框架。从中可以看出,引入的 GLL 位于层 2 和层 3 之间的 2.5 层,其主要的功能是完成对 IP 层与接入链路层之间以

① 973 计划(2007CB310601)和国家科技重大专项(2009ZX03004-001)资助项目。

② 男,1982 年生,博士;研究方向:异构无线网络的融合及资源管理,无线通信网,多媒体通信;联系人, E-mail: cuiyang_0305@163.com (收稿日期:2010-07-02)

及不同的接入链路层之间的协议转换和执行多接入分组调度。为了获得更多的多用户分集和多接入分集收益,本文采用将用户数据包选择和接入链路分配合成一步的多接入分组包调度方法。整个通用链路层主要由用户 QoS 分类器,先进先出(FIFO)队列和多接入分组调度器这三部分构成。从 IP 层到来的实时业务数据包流首先按照 QoS 进行分类并移入到相应的 FIFO 队列中,然后分组调度器在每个调度周期内依据调度准则进行用户分组包的选择和 RA 的分配。其大致过程如图 1 所示。本文为了分析方便,假设所有用户均只产生一个实时分组业务,对于多个业务的情况本文所提的算法也适用。

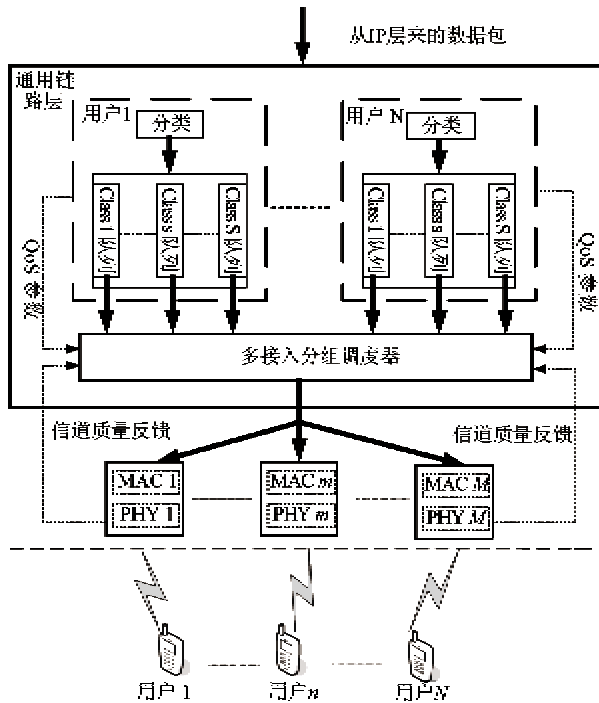


图 1 实时业务多接入调度框架图

2 多接入分组调度模型和 QU 算法

2.1 调度效用函数

为了表示执行一次多接入分组调度（即选择某一用户的业务分组包并将其分配到某一个无线接入技术的 RA 上）在保证用户 QoS 和避免资源浪费方面带来多大的收益,本文利用效用函数来表示。对于实时分组业务来说,时延是一个重要的 QoS 参数。当等待发送的分组包的实时等待时间超过其延时门限 D 时,在发送端将会被丢弃。于是对于调度效用函数 $U_j(t)$ 来说,其值应该随着实际时延 $d_i(t)$ 的函数 $D_i(d_i(t))$ 增长而单调增长,换句话说,分组

包的实时时延 $d_i(t)$ 越大,其对应的调度效用函数值就越大,该包就越应该被调度。同时当分组包的实时时延接近门限 D_i 时, $D_i(d_i(t))$ 的值就急剧地增加。满足上述性质的 $D_i(d_i(t))$ 函数表达为^[8]

$$D_i(d_i(t)) = e^{d_i(t)/D_i} \quad (1)$$

其中 i 和 j 分别表示第 i 个用户和第 j 个 RA。根据在保障实时业务 QoS 的基础上避免资源浪费的这一原则,调度器应该将分组包尽量地调度到能够为其提供更大数据速率的 RA 上,同时还要避免将其调度到那些为其提供的数据速率超过保证 QoS 所需速率的 RA 上。为了清楚地说明这一问题,首先定义一些参数: $R^{\min}$ 为最小目标数据速率,也就是为了保证每个队列头的分组包不超时所需的最小速率,定义它为实时业务在每次调度时所需的最小速率,其计算公式为

$$R^{\min}_i(t) = \frac{USD_i}{D_i - d_i(t)} \quad (2)$$

其中 USD_i 是队列头分组包剩余没有被传输的比特数。如果该包是完整的数据包时, USD_i 为零。对于无线通信系统中的实时分组业务来说,由于链路的时变性难以对所有分组给出一个确定的时延保证,但可以从统计的角度给出 QoS 保证。文献[9]提出了一种实时业务时延要求的定义方式:

$$P\{d_i > D_i\} \leq \delta_i \quad (3)$$

式(3)的实际意义就是使得在队列中的分组包超时的概率不得高于 δ_i 。针对从统计角度给予实时分组业务 QoS 保证的问题,目前已经有大量不同的研究方法,其中有效带宽^[10]理论便是一种有效的解决方法。有效带宽最早应用于有线分组网络的研究中,但在无线网络中有效带宽也可以用来分析实时业务 QoS 的统计要求。假定在发送端,为每个实时业务分配一个缓存队列,其队列长度(单位为比特)等于业务最大时延与业务产生的平均数据速率的乘积。因此,式(3)中表示不能保证实时业务 QoS 的概率也就等于队列溢出的概率。

假定实时业务是一个静态源,那么确保式(3)成立所需的最小带宽 E_i 为

$$E_i = m_i + \frac{\rho_i}{2B_i} \log\left(\frac{1}{\delta_i}\right) \quad (4)$$

其中 m_i 代表第 i 个实时业务产生的平均速率; δ_i 为溢出概率; B_i 为缓存队列的长度; ρ_i 为弥散系数:

$$\rho_i = \lim_{t \rightarrow \infty} \frac{1}{t} E\{X^2[0, t]\} \quad (5)$$

其中 $X[0, t]$ 为时间在 0 到 t 内的业务总量。用

ON-OFF 模型来建模 VoIP 业务^[10],于是有效带宽为

$$E_i = \frac{\lambda_i x_i}{\lambda_i + \mu_i} + \frac{\lambda_i \mu_i x_i^2}{B_i(\lambda_i + \mu_i)^3} \log\left(\frac{1}{\delta_i}\right) \quad (6)$$

其中 x_i 为峰值速率; $1/\mu_i$ 和 $1/\lambda_i$ 分别为激活和静默时间的均值。定义 $R_{ij}(t)$ 为在 t 时刻第 j 个 RA 为第 i 个用户所提供的实际速率; $R_{\max i}(t)$ 为在 t 时刻第 i 个用户可以获得的最大链路速率; $\bar{R}_i(t)$ 为第 i 个用户在某一周期 t_c 内可以获得的实际平均速率,其表达式为^[11]

$$\bar{R}_i(t) = \left(1 - \frac{1}{t_c}\right) \cdot \bar{R}_i(t-1) + \frac{1}{t_c} \cdot R_{ij}(t) \cdot A_{ij}(t) \quad (7)$$

其中

$$R_{ij}(t) = \min\left\{r_{ij}(t), \frac{q_i(t)}{\Delta t_i}\right\} \quad (8)$$

A_{ij} 为无线接入链路的分配指示索引,当第 j 个 RA 被分配给第 i 个用户时, $A_{ij}(t) = 1$, 否则 $A_{ij}(t) = 0$ 。

$r_{ij}(t)$ 为 RA 为用户 i 所提供的峰值速率, q_i 为用户队列中所有的数据包之和的大小, Δt_i 为调度时隙的大小。为了从统计的角度给实时分组业务提供 QoS 保障,应该保证在一段时间内的平均速率 $\bar{R}_i(t)$ 等于有效带宽 E_i , 于是定义 $R_{Ei}(t)$ 为在 t 时刻使得平均速率等于有效带宽时所需的数据速率,称其为有效速率。当用户的分组包被调度到为其提供的速率超过 R_{Ei} 的 RA 上时,将被视为资源浪费,不利于频谱效率的提高。调度效用函数 $U_{ij}(t)$ 随着以 R_{ij} 作为自变量的权重函数 $RA_{ij}(R_{ij})$ 增长而增长,基于上面的分析, $RA_{ij}(R_{ij})$ 的定义如下:

情况 1: $R_{\min i} < R_{Ei} < R_{\max i}$

$$RA_{ij}(R_{ij}) = \frac{1}{1 + e^{-a_i(R_{ij}-b_i)}} \quad (9)$$

其中 $a = \frac{2\ln 10}{R_{Ei} - R_{\min i}}, b = \frac{2}{R_{Ei} - R_{\min i}}$ 。

情况 2: $R_{\max i} \leq R_{\min i} \leq R_{Ei}$

$$\begin{aligned} RA_{ij}(R_{ij}) &= 1 \quad (R_{ij} = R_{\max i}) \\ RA_{ij}(R_{ij}) &= 0 \quad (R_{ij} \text{ 为其他值}) \end{aligned} \quad (10)$$

情况 3: $R_{\min i} < R_{\max i} < R_{Ei}$

$$RA_{ij}(R_{ij}) = \frac{1}{1 + e^{-a_i(R_{ij}-b_i)}} \quad (11)$$

其中 $a = \frac{2\ln 10}{R_{\max i} - R_{\min i}}, b = \frac{2}{R_{\max i} - R_{\min i}}$ 。

第一种情况表明用户 i 的 R_{Ei} 和 $R_{\min i}$ 均可以得到满足,从图 2(a) 的函数曲线中可以清楚地看出当 R_{ij} 的值小于 $R_{\min i}$ 时, $RA_{ij}(R_{ij})$ 几乎为零;当 R_{ij} 大于 $R_{\min i}$ 时, $RA_{ij}(R_{ij})$ 随着 R_{ij} 急剧地增长,直到 R_{ij} 等于

R_{Ei} 为止。当 R_{ij} 大于 R_{Ei} 以后, $RA_{ij}(R_{ij})$ 几乎不再随着 R_{ij} 增长而显著增长,而是趋于平缓,其值接近于 1。由于调度器总是以最大化效用函数为目标,这样设计效用函数可以避免为那些已经得到 QoS 保障的用户分配过多的资源,进而在相同的资源下能够为更多的实时业务用户提供有 QoS 保障的服务。因此在用户数量较多的情况下,可以降低丢包率,提高整个系统的有效吞吐量(频谱效率)。在情况 2 中, R_{Ei} 和 $R_{\min i}$ 不能得到满足,或是 $R_{\min i}$ 和 R_{Ei} 等于为它提供最大速率 $R_{\max i}$, 为了尽量保证业务的 QoS,它只应该被调度到能够为它提供 $R_{\max i}$ 的 RA 上,于是在这个情况上,只有当 R_{ij} 等于 $R_{\max i}$ 时, $RA_{ij}(R_{ij})$ 为 1;其余的情况, $RA_{ij}(R_{ij})$ 均为 0。对于第三种情况,用户 i 的最大速率 $R_{\max i}$ 小于 R_{Ei} , 为了保证实时业务的 QoS,分组包应该被调度到能够为其提供的速率大于 $R_{\min i}$ 的 RA 上,图 2(b) 表明了 $RA_{ij}(R_{ij})$ 和 R_{ij} 的关系。

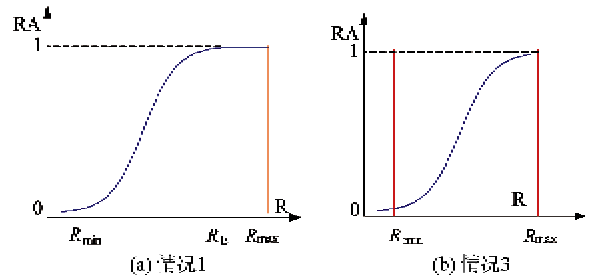


图 2 RA 函数曲线

基于上述的分析,调度效用函数被定义如下:

$$U_{ij}(t) = RA_{ij}(R_{ij}(t)) \cdot D_i(d_i(t)) \quad (12)$$

为了满足实时业务的 QoS 需求同时尽量提高频谱效率,调度器应该遵循最大化效用函数的原则来进行分组包的选择和接入链路的分配。

2.2 多接入调度模型

在构造完调度效用函数后,本文建立多接入分组调度的模型。假设每个用户只能同时产生一个实时业务,在队列中每个用户的实时业务数据包在调度的时候服从 FIFO 规则。于是在任一传输时间间隔内分组调度模型为

$$\max\left(\sum_{i=1}^N \sum_{j=1}^M U_{ij}(t) \cdot V_{ij}(t)\right) \quad (13)$$

$$\max\left(\sum_{i=1}^N \sum_{j=1}^M \frac{R_{ij}(t)}{\bar{R}_i(t)} \cdot V_{ij}(t)\right) \quad (14)$$

$$\text{s. t. } \sum_{j=1}^M V_{ij}(t) = 1 \quad (15)$$

其中式(13)被用来保证实时业务的 QoS 和提高频谱效率。式(14)是用来保证每个用户的公平性,避免一些信道条件较差的用户始终无法得到调度。在系统中每个 RA 在任意一传输时间间隔内只能被分给一个用户使用, N 和 M 分别为用户和可以利用的 RA 数量。很显然所建立的这个调度模型是一个 NP 问题^[12]。通常 NP 问题可以通过一些优化算法来解决,例如遗传算法^[13]和神经网络^[14]。对于算法的设计,执行时间这一因素不能不被考虑,因为算法的执行时间必须远小于调度周期。本文采用的调度周期等于传输时间间隔,是毫秒级的,因此需要一个快速、有效地求解调度模型的算法,而 Hopfield 神经网络(HNN)恰恰是快速地求解 NP 问题是一个有效的工具。除此之外,其易于硬件实现,而且通过硬件实现后的 HNN 求解所需的时间是微秒级的,其硬件实现的结构图如图 3 所示。基于上述分析,本文提出了基于 HNN 的分组调度算法。

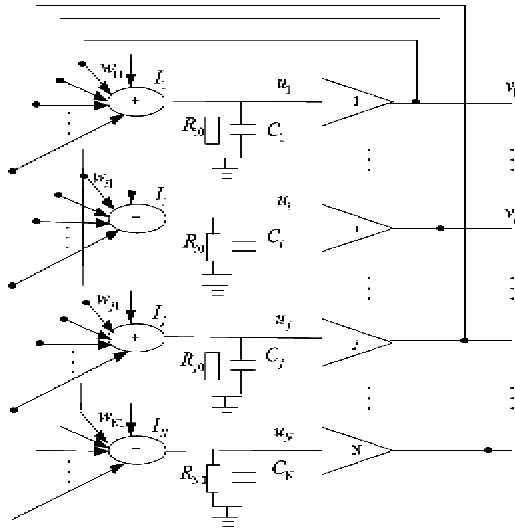


图 3 连续型 Hopfield 神经网络

2.3 Hopfield 神经网络和 QU 调度算法

HNN 是互联型网络的一种,它引入类似于 Lyapunov 函数的能量函数概念,把神经网络的拓扑结构与所求问题相对应,并将其转换为神经动力系统的演化问题。其演变过程是一个非线性动力学系统,可以用一组非线性差分方程描述(离散型)或是微分方程(连续型)来描述。系统的稳定性可用所谓的“能量函数”进行分析。在满足条件的情况下,某种“能量函数”的能量在网络运行过程中不断地减少,最后到达区域稳定的平衡状态。如果把系统的稳定点视为一个能量函数的极小点,而把能量函

数视为一个优化问题的目标函数,那么从初态朝这个稳定点的演变过程就是一个求解该优化问题的过程。因此,HNN 的演变过程就是求解优化问题的过程。实际上,它的解决并不需要真的去计算,而是通过构成反馈神经网络,适当地设计其连接权和输入就可以达到这个目的。

本文采用连续型的 HNN,它是由一系列相互连接的神经元组成,这些神经元动态地改变它们的输出直到整个网络到达平衡点。Hopfield 应经证明了利用能量函数 E 表示输出的动态变化,同时可以通过寻求能量函数 E 的局部最小值来达到获取平衡点的目的^[15]。HNN 的动态过程如

$$\frac{dU_{ij}}{dt} = -\frac{U_{ij}}{\tau} - \frac{\partial E}{\partial V_{ij}} \quad (16)$$

式所示^[16]。其中 U_{ij} 和 V_{ij} 分别是神经元 (i,j) 的输入和输出, τ 是一个时间常数。本文用 Sigmoid 函数作为神经元的激活函数,其表达式为

$$V_{ij} = f(U_{ij}) = \frac{1}{1 + e^{-\alpha_{ij} U_{ij}}} \quad (17)$$

其中 α_{ij} 为神经元的 shape 参数, $U_{ij} \in [-\infty, +\infty]$, $V_{ij} \in [0, 1]$ 。为了解决上面提出的问题,本文使用一个带有 $N \cdot M$ 个神经元的 2-D HNN。它的能量函数 E 被定义为

$$\begin{aligned} E = & -\mu_1 \sum_{i=1}^N \sum_{j=1}^M U_{ij}(t) \cdot V_{ij}(t) \\ & -\mu_2 \sum_{i=1}^N \sum_{j=1}^M \frac{R_{ij}(t)}{R_i(t)} \cdot V_{ij}(t) \\ & + \frac{\mu_3}{2} \sum_{i=1}^N \sum_{j=1}^M V_{ij}(t) (1 - V_{ij}(t)) \\ & + \frac{\mu_4}{2} \sum_{i=1}^N (1 - \sum_{j=1}^M V_{ij}(t))^2 \end{aligned} \quad (18)$$

式中第 3 项用来保证神经元可以快速地收敛于一个正确和稳定的状态,即 1 或 0,当 $V_{ij} = 1$ 时表示用户 i 的数据包被调度到 RA_j 上; $V_{ij} = 0$ 表示用户 i 的数据包没有被调度到 RA_j 上。第 4 项用来表示在任意一个传输时间间隔内,一个 RA 只能被分配给一个用户,但一个用户在一个传输时间间隔内可以使用多个 RA。根据 Euler 规则,对式(16)进行如下的迭代:

$$U_{ij}(t + \Delta t) = U_{ij}(t) + \Delta t \left\{ -U_{ij}(t) - \frac{\partial E}{\partial V_{ij}} \right\} \quad (19)$$

其中 Δt 是神经元输出更新的时间间隔,能量函数的梯度可以按照下式进行计算:

$$\frac{\partial E}{\partial V_{ij}} = -\mu_1 U_{ij}(t) - \mu_2 \frac{R_{ij}(t)}{R_i(t)} + \frac{\mu_3}{2}(1 - 2V_{ij}(t)) - \mu_4(1 - \sum_{s=1}^M V_{is}(t)) \quad (20)$$

3 仿真及结果

3.1 仿真参数

本文在仿真的场景中,只考虑将数据下行传输给 N 个移动用户,分别属于两个不同无线接入技术的 M 个可利用 RA。建立的仿真网络包含 7 个采用无线接入技术 1 的小区 and 7 个采用无线接入技术 2 的小区,这两类小区的半径分别为 3km 和 1.5km,每个小区下行的总的可用带宽为 250kHz。其具体结构如图 4 所示。所有的无线接入技术均采用时分复用接入(TDMA)方式,它们的传输时间间隔均为 1ms,并将其作为调度时间间隔。在仿真中,RA 之间的信道衰落是不同的,路径损耗和阴影衰落采用郊外模型,多径衰落采用修改的 Jakes 模型。在仿真的初始阶段,所有的用户一致地分布于图 4 中的深色区域,同时为了避免边界效应使用了 wrap-around 技术。终端的移动速率是 3km/h,这样的话可以保证在每个传输时间间隔内,信道的条件不变。

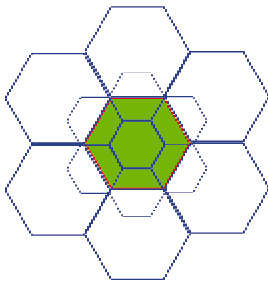


图 4 仿真场景

所有无线接入链路的物理层均采用单载波编码调制,具体的调制编码方案选择策略及其获得的频谱效率如表 1 所示^[17]。

表 1 调制编码方案选择策略

信噪比 (dB)	调制编码方案	频谱效率(bit/s/Hz)
$-\infty < \gamma(t) \leq 1.69$	QPSK 1/3 rate	2/3
$1.69 < \gamma(t) \leq 4.82$	QPSK 1/2 rate	1
$4.82 < \gamma(t) \leq 7.21$	16-QAM 1/3 rate	4/3
$7.21 < \gamma(t) \leq 11.8$	16-QAM 1/2 rate	2
$11.8 < \gamma(t) \leq \infty$	64-QAM 1/2 rate	3

3.2 业务模型

本文选择的实时业务是平均速率为 64kbps 的 VoIP,它可以用 interrupt 决策过程(ON 和 OFF 的持续时间服从的指数分布的 ON-OFF 模型)进行建模。VoIP 最大允许的时延和丢包率分别是 60ms 和 5%。为了评估所提的 QU 算法对于实时分组业务的性能,本文将 QU、PLR 和 M-LWDF 算法同时进行仿真,并在频谱效率、丢包率、平均时延以及时延抖动这 4 方面进行比较。

3.3 仿真结果

仿真结果如图 5~图 8 所示。在负载较轻时,QU、PLR 和 M-LWDF 这 3 种算法在频谱效率、丢包率和平均延时这 3 个方面的性能比较接近。但是随着用户数逐渐增大,这 3 种算法的性能出现了差异。这 3 种算法的频谱效率都随着用户数量的增加而增加,但当用户数量增大到 62 后,QU 算法的频谱效率要比其它两种算法要大,当用户数变为 68 时,这种优势最大,其频谱效率大约提高了 5%。主要原因是当其它条件相等时,例如实际时延和公平性,这两

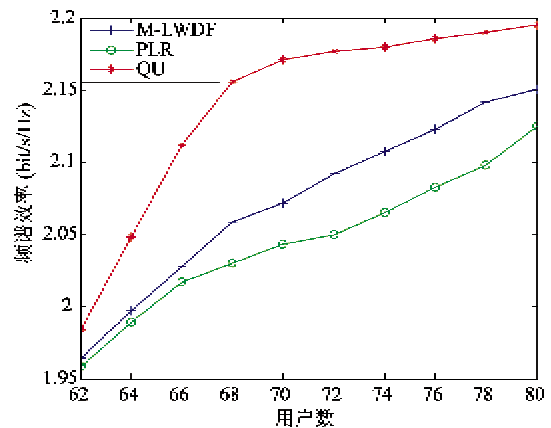


图 5 三种算法的频谱效率曲线比较

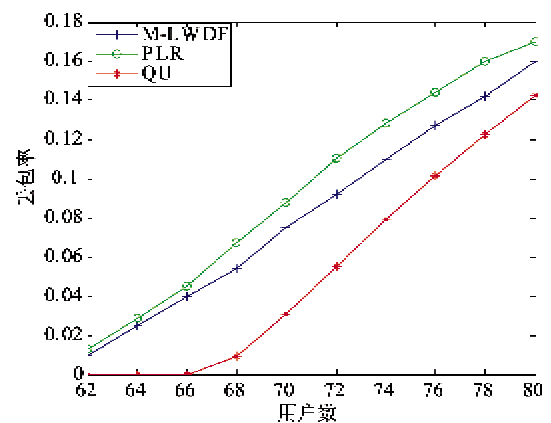


图 6 三种算法的丢包率曲线比较

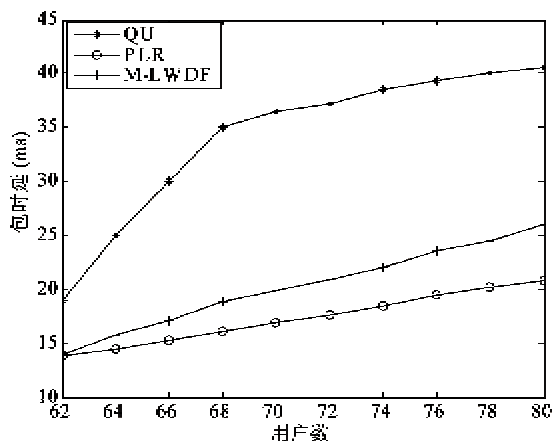


图7 包的平均延时曲线比较

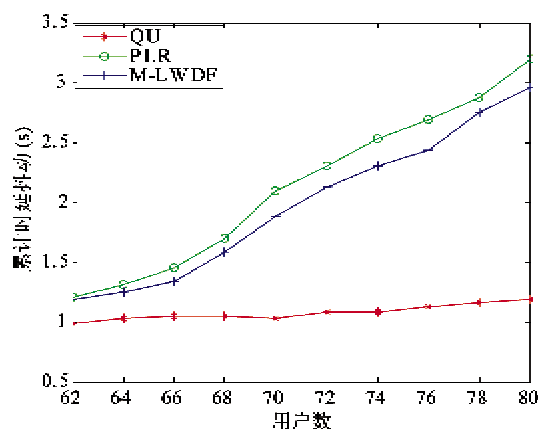


图8 累计分组时延抖动比较

个算法倾向于将用户包调度给能够为其提供最大数据速率的 RA 上,因此在一定程度上可以提高频谱效率。但是 QU 算法将怎样避免频谱资源的浪费考虑进算法中,其具体体现在式(9) - (11),由于在系统的负载达到一定程度后采用这种机制的 QU 算法相比于其它两种算法可以提供更低的丢包率,进而提高了整个系统的有效吞吐量,因此也就提高了整个系统的频谱效率。同理也可以解释为 QU 算法可在满足 QoS 要求的情况下能够为更多的用户提供服务。

上述同样的结论也可以从图7中得出。在用户数未超过66时,QU算法所提供的丢包率几乎为零,在所有不同的用户数情况下QU算法相比这两个经典算法能够提供更低的丢包率。

在平均延时方面,QU算法在负载变大以后平均延时要比其他两种算法要高,这是由于QU算法采用避免为满足QoS需求后的实时业务提供过多的资源所造成的,但是值得注意的是,QU算法虽然在延时方面比其他两种算法大,但是仍没有超过实

时分组业务所允许的最大时延。

图8为任意选择一个用户在不同用户数量条件下的累计时延抖动。从该图中可以清楚地看出在不同的负载情况下,QU算法的时延抖动累计量均为最低,而且并没有随着负载的增大而显著增加。这主要还是要归功于式(9)-(11)中引入的资源避免机制,它在用户信道条件好的时候并不为其分配过多的资源,同时也保证了在信道条件不好时为其分配足够的资源。综上分析,本文所提的多接入分组调度算法QU能够为实时业务提供更好的QoS保障。

4 结论

本文引入通用链路层作为异构无线网络融合的解决方案,并针对在通用链路层上运行的实时分组业务提出了具有QoS保障的多接入分组调度算法QU。该算法充分利用了多用户分集与多接入分集带来的增益,在保证实时分组业务QoS需求的同时提高了无线频谱资源的使用效率。与经典的M-LWDF和PIR算法相比,QU算法在满足时延需求的基础上,提高了频谱效率,降低了由于超时所造成的丢包率并减少了时延抖动。而且该算法为实时业务节约资源的这一特点也为通信系统能够容纳更多的非实时分组业务提供了可能。同时该算法使用了HNN对多接入分组调度模型进行优化求解,相对于其它的优化算法而言,具有易于硬件实现和能够快速求解的特点,因此便于在实际系统中应用。

参考文献

- [1] Yang Z. Heterogeneous networks convergence and cooperation. *ZTE Communications*, 2008,14(3):1-2
- [2] Sach J, Wienmnn H, Magnusson P, et al. A generic link layer in a beyond 3G multi-radio access architecture. In: *Proceedings of the 26th IEEE Conference on Communications, Circuits and Systems*, 2004 ICCAC, Helsinki, Finland, 2004. 447-451
- [3] Dimou K, Agliero R, Bortnik M, et al. Generic link layer: a solution for multi-radio transmission diversity in communication networks beyond 3G. In: *Proceedings of the 26th IEEE Vehicular Technology Conference*. Paris, France, 2005. 3. 1672-1676
- [4] IEEE Standard of Local and Metropolitan Area Networks-Part16: Air Interface for Fixed and Mobile Broadband Wireless Access System-Amendment2: Physical and Medi-

- um Access Control Layers for Combined Fixed and Mobile Operation in Licensed Bands, 2006
- [5] Mobile broadband: The global evolution of UTM/HSPA-3GPP release 7 and beyond. 3G Americas whitepaper, July, 2006. <http://whitepapers.techrepublic.com.com/abstract.aspx?docid:308653>
 - [6] Ramannan K, Stolyar A, Whiting P, et al. Providing quality of service over a shared wireless link. *IEEE Communications Magazine*, 2000, 39(2): 150-154
 - [7] Liu G, Zhang J H, Zhou B, et al. Scheduling performance of real time service in multiuser OFDM system. In: Proceedings of the IEEE Conference on Wireless Communications, Networking and Mobile Computing, Shanghai, China, 2007. 504-507
 - [8] Sang A, Wang X, Madihian M. Real-time QoS in enhanced 3G cellular packet systems of a shared downlink channel. *IEEE Transactions on Wireless Communications*, 2007, 6(5): 1803-1812
 - [9] Andrews M, Kumaran K, Ramanan K, et al. Providing quality of service over a shared wireless link. *IEEE Communication Magazine*, 2001, 39(2): 150-154
 - [10] Courcoubetis C, Fouskas G, Weber R. On the performance of an effective bandwidth Equation. In: Proceedings of the 14th International Teletraffic Conference, London, UK, 1994, 5. 6-10
 - [11] Jalali A, Padovani R, Pankai. R. Data throughput of CDMA HDR a high efficiency-high data rate personal communication wireless system. In: Proceedings of the IEEE Vehicular Technology Conference, Tokyo, Japan, 2000, 3. 1854-1858
 - [12] Kechrotis G I, Manolakos E S. Hopfield neural network implementation of optimal CDMA multiuser. *IEEE Transactions on Neural Network*, 1996, 7:131-141
 - [13] Wang J, Huang J, Rao S Q, et al. An adaptive genetic algorithm for solving traveling salesman problem. In: Proceedings of the 4th International Conference on Intelligent Computing, HongKong, China, 2008. 182-189
 - [14] Ercsey-Ravasz M, Roska T, Neda Z. Cellular neural networks for NP-hard optimization. In: Proceedings of the 11th International Workshop on Cellular Neural Networks and Their Applications, Santiago, Chile, 2008. 52-56
 - [15] Abe S. Theories on the Hopfield neural networks. In: Proceedings of the International Joint Conference on Neural Networks, New York, USA, 1989, 1(18):557-564
 - [16] Almajano L, Pérez-Romero J. Packet scheduling algorithm for interactive and streaming service under QoS guarantee in a CDMA system. In: Proceedings of the IEEE Vehicular Technology Conference, Berlin, Germany, 2002, 3. 1657-1661
 - [17] Karimi H R, Koudouridis G P, Dimou K. On the spectral efficiency gains of switched multi-radio transmission diversity. In: Proceedings of the 25th Conference on Wireless Personal Multimedia Communications, Tokyo, Japan, 2005. 1177-1181

Multi-radio packet scheduling for real-time packet traffic on general link layer

Cui Yang, Xu Yubin, Xu Rongqing, Sha Xuejun, Ding Zhe

(Communication Research Center of Harbin Institute of Technology, Harbin 150080)

Abstract

In this paper a multi-radio packet scheduling algorithm, called the QU algorithm, is presented to guarantee the QoS for the RT (real-time) traffic on the general link layer. The main contributions of this study are as follows: Firstly, with the purpose of guaranteeing RT traffic' QoS and avoiding wasting radio resources, a scheduling-utility function was developed to represent the utility of executing a packet scheduling; Secondly, using this utility function, a multi-radio packet scheduling model was established and analyzed based on joint consideration of fairness; Thirdly, the Hopfield neural network was used to fast calculate an optimal solution of this scheduling model. The simulation results show that, compared with M-LWDF and PLR algorithms, the QU algorithm has the higher spectrum efficiency, lower packet loss ratio and can delay variation while meeting allowable average packet delay under high load.

Key words: general link layer (GLL), multi-radio packet scheduling, scheduling-utility, neural network