

# 以太网上实时业务服务质量保证策略的研究

黄周松 雷振明

(北京邮电大学信息工程学院 ATM 中心 北京 100876)

**摘要** 首先分析了以太网上相关技术的进展,并指出了以太网上保证实时业务服务质量的迫切性,然后提出了以太网上保证实时业务服务质量的策略,并给出了该策略的实现结构。仿真验证了我们提出的策略对实时业务服务质量的保证是有效的。

**关键词** 虚拟局域网,多生成树协议,服务质量

## Study of the Policy to Assure the Quality of Real-time Services on Ethernet

HUANG Zhou-Song LEI Zhen-Ming

(ATM Center of Information Engineering Institute, Beijing University of Posts and Telecommunication, Beijing 100876)

**Abstract** This paper first analyzes the progress of Ethernet technologies. And at the same time, the necessity of assuring the quality of real-time services on Ethernet is proposed. Then we propose a policy to assure the quality of real-time services on Ethernet and we bring forward the implementation framework. Simulation shows that the policy is efficient of assure the quality of real-time services.

**Keywords** Virtual LAN, Multiple STP, Quality of Service

## 1 引言

随着以太网技术的快速发展及千兆以太网标准<sup>[1]</sup>的提出,以太网技术已经能够满足高速网络的多种关键需求。因此,以太网技术的应用范围已经由局域网(LAN)扩展到城域网(MAN)<sup>[2]</sup>。同时,越来越多新业务的出现使得以太网承载的业务类型也不断增加,不仅要为传统的普通数据传输业务提供服务,还要传送实时的语音或图像数据。不同的用户和不同的应用,对网络服务质量有着不同的需求,这为以太网提出了挑战:要保证业务服务质量尤其是实时业务的服务质量。802.1p 机制<sup>[3]</sup>为以太网上支持与保证实时业务服务质量打下了基础,它为以太网上支持 QoS 走出了第一步,但它必须结合接纳控制技术才能从根本上解决以太网上保证实时业务服务质量这个难题。以太网上 VLAN 流量均衡技术也得到了一定的发展,VLAN 流量均衡能延缓瓶颈链路的过早出现,能在一定程度上改善业务的服务质量<sup>[2]</sup>,但是单纯基于 VLAN 流量均衡并不能保证以太网上实时业务的服务质量。为了解决以太网上实时业务服务质量保证这个难题,本文提出了一种基于测量的以太网上实时业务服务质量保证的策略,该策略结合了 VLAN 流量均衡、802.1p 机制、接纳控制与时延测量的技术。本文第 2 节描述了以太网上实时业务服务质量保证的策略;第 3 节给出了该策略的实现结构;第 4 节对该策略的性能进行了仿真;最后对全文进行了总结。

## 2 以太网上实时业务服务质量保证的策略

本节提出以太网上实时业务服务质量保证的策略,它包括四个部分,分别是:根据 802.1p 划分优先级、VLAN 划分的原则、VLAN 里实时业务服务质量的监测以及接纳与路由

控制。

根据 802.1p 划分优先级:为了减小非实时业务对实时业务性能的影响,根据 802.1p 给业务设置优先级,并选择进行优先级排队。

VLAN 划分的原则:每当一个实时业务流经过的路径上没有运行一个 VLAN 的话,则在这条路径里分配一个 VLAN;也就是说,VLAN 划分是针对路径的,而不是针对逐流的,这有利于节省 VLAN 资源,减少 VLAN 划分与配置的次数;并在不同的生成树 STP 里分配连接到所有接入网关的 VLAN 连接,为接纳与路由控制打下基础。

VLAN 里实时业务服务质量的监测:以 802.1D<sup>[4]</sup>标准规定的话音业务的时延和时延抖动均不超过 10ms,视频业务的时延和时延抖动均不超过 100ms 为监测目标。我们测量每个 VLAN 里每个实时业务流的报文经历的最大时延值  $D_{\max}$  与最小时延值  $D_{\min}$ 。时延的测量模型是:

$$D_i = (S_i - C_i)$$

其中  $S_i$  指的是报文到达网络出口处的接入网关的时间,  $C_i$  指的是实时业务报文中 RTP 报头域中的时戳;  $D_i$  指的是报文经历的网络传输时延。最大时延值  $D_{\max}$  就是  $D_i$  的最大值;最小时延值  $D_{\min}$  指的是  $D_i$  的最小值,于是我们得到某个流的最大时延抖动  $(D_{\max} - D_{\min})$ , 记为  $\Delta$ , 于是有  $\Delta = (D_{\max} - D_{\min})$ 。由于以太网上流量的自相似特性,它一般会导致报文经历时延的剧烈变化,而当  $\Delta$  非常小时,这意味着网络时延变化是非常平稳的,我们认为网络此时的性能非常好;而且当  $\Delta$  越小,一般就可以认为网络的性能越好。根据最大时延抖动  $\Delta$  来监测网络性能的另一个好处就是我们不需要校正源端与接入网关之间存在的时间偏差,这是因为时间偏差在我们做  $(D_{\max} - D_{\min})$  这个计算时消去了。

为了减少偶然性,在某个 VLAN 里的多个流中,我们选择最大的  $\Delta$  来评价该 VLAN 的性能。

接纳与路由控制:我们优先选择  $\Delta$  值最小的 VLAN 来接纳新的实时业务流。我们首先判断  $\Delta$  值是否小于一个较小的门限值  $\lambda_1$ ,如果是,接纳该请求;如果  $\Delta$  值大于一个较大的门限值  $\lambda_2$ ,则拒绝该请求;如果  $\Delta$  值在这两个门限之间,则需要判断该路径上理论对时延上限 ( $D_{\max}$ ) 是否超过 10ms 或者 100ms (对于语音业务则为 10ms,对于视频业务则为 100ms);若在 10ms 或者 100ms 之内,接纳该请求,否则将拒绝该请求。下面将给出理论时延上限的计算方法以及路由与接纳控制过程的伪代码。考虑到城域网节点之间距离相对较小,为简化分析,我们暂不考虑报文的链路传输时延。

对于语音业务,因为分配的优化级是 6<sup>[3]</sup>,考虑到网络控制流量一般比较小,我们忽略它对话音业务的影响。且由于语音业务速率基本固定,我们把它作为 CBR 类业务来处理。我们不妨设某个路径由  $n$  个链路组成,我们假设链路带宽均为  $r$ ,每条链路上分析有  $m_1, m_2, \dots, m_n$  个话音流新加入该路径;设链路上话音报文的最大长度为  $L$ ,链路上经历的报文的最大长度设为  $L_{\max}$ ,下面我们计算语音业务的理论时延上限  $D_{\max}$ 。

当某个话音流的报文到达第一个节点后,可能其它 ( $m_1 - 1$ ) 个话音流也有报文稍稍早于这个报文到达该交换节点,由于同一优先级的报文需要根据到达先后顺序进行转发,因此它需要等待其它 ( $m_1 - 1$ ) 个报文转发完毕后才能得到转发;而在转发该话音流报文之前节点可能正在转发一个长度为  $L_{\max}$  的低优先级报文,这将增加这个话音流报文的等待时延;然后,该报文得到转发,其转发时延为  $\frac{L}{r}$ 。因此该报文被第一个节点转发需要经历时延的最大值  $D_{\max 1}$  可以表达为:

$$D_{\max 1} \leq \sum_{i=1}^{m_1-1} \frac{L}{r} + \frac{L_{\max}}{r} + \frac{L}{r} = \sum_{i=1}^{m_1} \frac{L}{r} + \frac{L_{\max}}{r}$$

而该话音流报文还需从第 2 段链路传输到第  $n$  段链路,这段链路经历的时延最大为 ( $D_{\max} - D_{\max 1}$ );考虑到  $D_{\max}$  的上限为 10ms,因此从第 2 段链路传输到第  $n$  段链路的时延将小于 10ms。由于语音流的发送时间间隔一般是 20ms 或更大,因此在经历第 2 段至第  $n$  段链路时最多有 ( $m_2 + m_3 + \dots + m_n$ ) 个话音报文会对该报文造成进一步的延迟;考虑到该话音报文经历 ( $n-1$ ) 条链路的转发时延以及低优先级报文对它的影响,我们得到:

$$D_{\max} \leq D_{\max 1} + \sum_{i=2}^n \frac{m_i \times L}{r} + (n-1) \frac{L_{\max}}{r} + (n-1) \frac{L}{r}$$

结合  $D_{\max 1}$  的表达式,我们得到:

$$D_{\max} \leq \frac{\sum_{i=1}^n m_i \times L}{r} + \frac{n \times L_{\max}}{r} + (n-1) \frac{L}{r}$$

这就是话音报文经历一个路径后的理论时延上限,从上式可以看出它决定于路径上接纳的流的数目以及该路径的跳数,此外,它还受报文长度的影响。

对于视频业务,由于分配的优先级为 5<sup>[3]</sup>,它的服务质量除了受同优先级业务的影响,还受到优先级为 6 的话音业务的影响;我们不妨设某条路径由  $n$  个链路组成,链路带宽均为  $r$ ,每条链路上分别有  $m'_1, m'_2, \dots, m'_n$  个视频流新加入该路径,设链路上视频流的最大报文长度为  $L'$ ,设链路上最大报文长度为  $L_{\max}$ ;设每个视频流经过漏桶整形,记第  $n$  条链路上新加入的第  $m'_n$  个视频流的平均速率为  $\rho_{m'_n}$ ,记第  $n$  条链路上新加入

的第  $m'_n$  个视频流的最大突发流量为  $b_{m'_n}$ ;下面我们将计算视频业务的理论时延上限  $D'_{\max}$ 。

当某个视频流的报文到达第一个节点后,可能其它 ( $m'_1 - 1$ ) 视频流也有报文稍稍早于这个报文到达;此外,这个流中可能还有报文先于它到达但还没有得到转发,因此它需要等待  $m'_1$  流中先于它到达的报文转发完毕后才能得到转发;而在转发该视频流报文之前节点可能正在转发一个长度为  $L_{\max}$  的低优先级报文,这将增加该视频流报文的等待时延;然后,该报文得到转发,其转发时延为  $\frac{L'}{r}$ 。因此该报文经过第一个节点需要经历时延的最大值  $D'_{\max 1}$  可以表达为:

$$D'_{\max 1} \leq \sum_{i=1}^{m'_1} \frac{b_i}{r} + \frac{L_{\max}}{r} + \frac{L'}{r}$$

而该视频流报文还需从第 2 段链路传输到第  $n$  段链路,这段链路经历的时延最大为 ( $D'_{\max} - D'_{\max 1}$ );因此在经历第 2

段至第  $n$  段链路时最多有 [ $\sum_{i=2}^n \sum_{j=1}^{m'_i} b_{ij} + (D'_{\max} - D'_{\max 1}) \times \sum_{i=2}^n \sum_{j=1}^{m'_i} \rho_{ij}$ ] 这么多流量会对该视频流报文造成进一步的延迟;考虑到该视频流报文经历 ( $n-1$ ) 条链路的转发时延、高优先话音报文以及低优先级报文对它的影响,我们得到:

$$D'_{\max} < \frac{\sum_{i=1}^n \sum_{j=1}^{m'_i} b_{ij}}{r - \sum_{i=2}^n \sum_{j=1}^{m'_i} \rho_{ij}} + \frac{n(L_{\max} + L')}{r} + D_{\max}$$

这就是视频流报文经历一个路径后的理论时延上限,从上式可以看出它决定于路径上接纳的流的突发流量、流的平均速率以及该路径的跳数,此外,它还受报文长度的影响。

下面给出接纳与路由控制过程的伪代码,设两个门限分别为  $\lambda_1$  与  $\lambda_2$ ,对于语音业务应满足  $0 < \lambda_1 < \lambda_2 < 10\text{ms}$ ;对于视频业务则应满足  $0 < \lambda_1 < \lambda_2 < 100\text{ms}$ ;对应的伪代码为:

```
// 优先选择  $\Delta$  最小的 VLAN 进行接纳控制
if ( $\Delta \leq \lambda_1$ )
    Accepted;
else if ( $\Delta \leq \lambda_2$ )
    if (( $D_{\max} < 10\text{ms}$ ) && ( $D'_{\max} < 100\text{ms}$ ))
        Accepted;
    else
        rejected;
end if;
else
    rejected;
end if;
```

本文没有采取文[5,6]中提到的 WSP (Widest-Shortest Path) 或 SWP (Shortest-Widest Path) 这两种经典的选路算法来进行选路;这是因为这两种算法设计的主要目的是为流量均衡服务的,它们有利于延缓瓶颈链路的过早出现,更适合于速率保证这类带宽租用业务的选路;但是它们在选路时没有考虑时延性能需求,也没有考虑流的突发对实时业务性能的影响,因此它们不适合本文选路设计的要求。

### 3 以太网上实时业务服务质量保证策略的实现结构

在确定了相关的策略后,下一个关键问题是如何建立一个合适的实现结构以利用该策略实现以太网上实时业务服务质量的保证。我们已经知道,以太网上流量均衡的实质是确定 VLAN 与 MSTP 的生成树之间的映射关系,而在 MSTP 协议体系中采用的集中式控制方法,映射关系的确定是由域中

唯一的主节点负责。因此,我们设计的实现结构也采用集中式结构,通过中心管理设备(Center Management Equipment, CME)集中实现接纳控制、选路功能。图1给出了以太网上实时业务服务质量保证策略的实现结构。该结构的核心是实现接纳控制与选路的中心管理设备CME。CME将预先规划好的STP拓扑规划与VLAN分配的信息配置到各个交换节点里,然后对实时业务流的请求进行接纳与路由控制。

用户的实时业务流量接入的基本过程为:位于局域网(LAN)中的用户向AG(Access Gateway,接入网关)发出连接请求;AG若发现该用户请求还没有得到认证与选路;则向CME转发该请求信息;CME接收到用户请求信息后,根据接纳与路由控制策略确定是否接纳该请求;如果接纳,为该用户业务流分配合适的路由并将其对应VLAN号反馈到转发该连接请求的AG。CME可以是某一个预先配置好的AG,也可以是一个独立的设备。对于非实时业务,在通过AAA认证后,AG将它接纳在CME预先配置好的VLAN里,而不需要CME进行动态干预。这有利于减小CME与AG之间交互的流量。

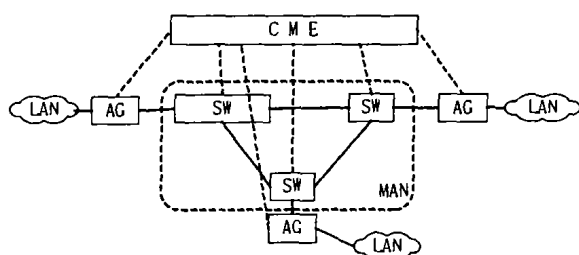


图1 实现结构

#### 4 性能仿真

为了验证本文提出的实时业务服务质量保证策略的性能,我们在Network Simulator2<sup>[6]</sup>里仿真了根据本文提出的策略与根据SWP或WSP算法进行接纳控制后实时业务的获得时延性能,并对它们进行了比较。以10ms作为以太网上传输时延的上限,并以此为目标;仿真时我们将参数 $\lambda_1$ 、 $\lambda_2$ 分别设置为7ms与9ms,仿真拓扑如图2所示。

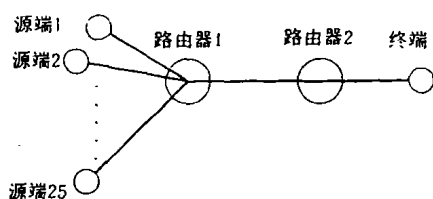


图2 仿真的网络拓扑图

图2中各段链路均为双向对称链路,各段链路的带宽均为1Mbps,每段链路的链路时延均为设为0.1ms。路由器采取Drop-tail调度机制,且其buffer大小设置为1000。源端1至25均产生到终端的流量请求,而终端是个Agent/Null类型的终端,在它里面加入了时延测量与接纳控制的功能。各个流随机申请接入,若允许接入,源端均每隔20ms产生一个报文,仿真中报文长度均为92字节。仿真持续了80s,仿真结果总体如下:

当接入了18个流后,终端测量到最大时延抖动超过 $\lambda_1$ ,并且此时接入流的理论时延上限已经超过10ms,因此,不再接纳别的流的请求。根据仿真结果的跟踪文件,我们整理出了

这18个流接入后的最大时延,以评价是否达到预期目标。这18个流经历的网络传输时延的最大值Max见表1,它们均以ms为单位。

表1 根据本文表提出的策略进行接纳控制时各流的传输时延的最大值

| 流号  | 1    | 2    | 3    | 4    | 5    | 6    |
|-----|------|------|------|------|------|------|
| Max | 7.66 | 9.13 | 7.66 | 9.13 | 8.39 | 7.66 |
| 流号  | 7    | 8    | 9    | 10   | 11   | 12   |
| Max | 8.39 | 7.66 | 9.13 | 7.66 | 9.13 | 8.39 |
| 流号  | 13   | 14   | 15   | 16   | 17   | 18   |
| Max | 8.39 | 7.66 | 9.13 | 7.66 | 9.86 | 8.39 |

从表1的数据我们可以看出本文提出的策略对最大时延是有保证的,接入的18个流使得链路的利用率达到66.24%,它的最大时延不超过9.86ms。下面将它与只根据理论时延上限或者流量均衡的思想进行接纳控制时带来的链路利用率与时延性能的差异进行比较。如果只根据理论时延上限进行接入,则最多可以接入13个流,可以计算出其链路利用率仅为47.84%;如果根据SWP或WSP这类流量均衡算法进行接入,则25个流均可以接入,此时链路利用率为92%,我们也对这种情况进行了仿真,为与表1进行比较,我们在表2中给出了其中18个流获得了最大传输时延。

表2 根据流量均衡算法进行接纳控制时各流的传输时延的最大值

| 流号  | 1     | 2     | 3    | 4    | 5    | 6    |
|-----|-------|-------|------|------|------|------|
| Max | 10.17 | 12.81 | 9.43 | 12.8 | 12.1 | 9.43 |
| 流号  | 7     | 8     | 9    | 10   | 11   | 12   |
| Max | 12.1  | 8.71  | 12.8 | 9.43 | 12.8 | 12.1 |
| 流号  | 13    | 14    | 15   | 16   | 17   | 18   |
| Max | 11.3  | 9.43  | 11.3 | 8.69 | 9.86 | 8.7  |

从仿真结果我们可以看出若根据SWP或WSP进行接纳控制的话,我们可以获得高的链路利用率,但是它对时延性能没有保证,从表2可以看出多个流经历的传输时延的最大值超过了10ms的性能要求;而根据本文提出的策略进行接纳控制,虽然在链路利用率上要低一些,但是在时延性能上是能够保证的。

**总结** 本文提出了一种以太网上保证实时业务服务质量的策略,并给出了它的实现结构。该策略根据对时延上限的理论分析与时延的实际测量结果进行时业务的接触控制,我们通过仿真验证了该策略能有效保证实时业务服务质量。

#### 参考文献

- 1 Draft Standard IEEE 802.3z, Media Access Control (MAC) Parameters, Physical layer, Repeater and Management Parameters for 1000Mb/s Operation. 1998
- 2 刘军. 以太网流量工程研究: [北京邮电大学博士学位论文]. 2003
- 3 IEEE Standard 802.1p incorporated. Media Access Control (MAC) Bridges; Revision. IEEE Standard for Local and Metropolitan area Networks, 1997
- 4 IEEE Standard 802.1D. Media Access control (MAC) Bridges. IEEE Standard for Local and Metropolitan Area Networks, 1998
- 5 柴洪杰. 以太网及相关技术的研究: [北京邮电大学博士学位论文]. 2002
- 6 Ma Qingming, Steenkiste P. On Path selection for Traffic with Bandwidth Guarantees. IEEE ICNP 97, 1997
- 7 NS2. Network simulator, Version 2. home page. <http://www.isi.edu/nsnam/ns>