

NEAR EAST UNIVERSITY

Faculty of Engineering

Department of computer Engineering

ATM NETWORKING

**Graduation Project
Com- 400**

Student: Mohamed Alameldin (960773)

Supervisor: Prof. Dr. Fakhreddin Mamadov

Lefkoşa -2000

ATM NETWORKING

CONTENTS

Acknowledgement	50
Abstract	50
Chapter 1 Introduction To N	50
1.1 Science And Network	51
1.2 Issues And Challenges	51
1.2.1 Networking Is Critical	51
1.2.2 Peak Of Technology Life Cycle	52
1.2.3 Standarization	53
1.2.4 Past Failures And Successes	53
1.2.4.1 Requirements for Success	53
1.3 The Situation In The Telecommunica	54
World	56
1.4 B-ISDN	57
1.4.1 Definition Of Brodband Services	
1.4.1.2 What Is Broadband	
1.4.3 Characteristic Of Broadband Services	
1.4.4 Broadband Services	
1.5 Service Introduction	
1.6 ATM Forum	
Chapter 2 Asynchronous Transfer Mode (ATM)	13
2.1 What Is ATM	13
2.2 How Does ATM Work	14
2.3 Who Is Using ATM	14
2.4 Why Do We Need ATM	15
2.4.1 ATM Benefits-I	16
2.4.2 ATM Benefits-II	16
2.5 Problems Addressed By ATM	17
2.5.1 Too Many Networks	18
2.5.2 Too Many Services	18
2.6 ATM Solution: The ATM Cell	19
2.7 The ATM Cell Structure	20
2.7.1 The Cell	20
2.7.2 Connection-Oriented	21
2.7.3 The Cell Size	21
2.7.4 The Cell Format	22
2.7.5 Time-Division Multiplexing	22
2.7.6 ATM Multiplexing	23
2.7.7 ATM Bandwidth And Cell Allocation	23

Chapter 3	ATM Network Concepts and Architecture	25
3.1	ATM's Position In The OSI Reference Model	26
3.1.1	Router Based Networking	27
3.1.2	X.25 Packet Switched Network	28
3.1.3	Switching Versus Routing	28
3.1.4	Frame-Relay Network	29
3.2	B-ISDN Protocol Reference Model	30
3.2.1	ATM as MAC Layer Protocol	31
3.2.2	ATM As a Link Layer Protocol	32
3.2.3	ATM As a Network Layer Protocol	32
3.2.4	ATM As A Transport Layer Protocol	33
3.3	ATM Functions And Layers	33
3.4	ATM Layer Functions	35
3.5	ATM Adaptation Layer Functions (AAL)	36
3.6	ATM Cell Structure Details	37
3.6.1	Generic Flow Control (GFC)	38
3.6.2	Virtual Path Identifier/ Virtual Channel Identifier (VPI/VCI)	38
3.6.3	Payload Type (PT)/ Cell Loss Priority (CLP)/ Header Error Control (HEC)	38
3.7	Virtual Channels / Paths	39
3.7.1	Virtual Channel / Path Connection	41
3.8	ATM Congestion Control	41
3.8.1	Why The Need Of Congestion Control	41
3.8.2	Connection Parameters	43
3.8.2.1	Quality Of Service	43
3.8.2.2	Usage Parameters	43
3.8.2.3	Service Categories	44
3.9	What Is Expected From Congestion Control	46
3.9.1	Selection Criteria	46
3.9.2	Generic Functions	47
3.10	Connection Admission Control	48
3.10.1	Usage Parameter Control	48
3.10.2	Generic Cell Rate Algorithm	48
3.10.3	Priority Control	49
3.10.3.1	Traffic Shaping	49
3.10.3.2	Leaky Bucket Algorithm	49
3.10.3.3	Network Resource Management	50

3.10.3.4	Frame Discard	50
3.10.4	Feedback Control	50
3.10.5	ABR Flow Control	50
3.10.5.1	Some Early Debates	51
3.10.5.1.1	Open Loop Vs Closed Loop	51
3.10.5.1.2	Credit-Based Vs Rate Based	51
3.10.5.1.3	Binary Feedback Vs Explicit Feedback	52
3.10.5.1.4	Congestion Detection Queue Length Vs Queue Growth	53
3.11	RM-Cell Structure	53
3.11.1	Service Parameters	54
3.11.2	Source, Destination and Switch Behavior	56
3.12	Representative Schemes	57
Chapter 4	A survey Of Switching ATM Techniques	61
4.1	Switching Functions	62
4.1.1	User Plane	63
4.1.2	Control Plane	63
4.1.3	Management Plane	63
4.1.4	Traffic Control Functions	64
4.2	A Generic Switching Architecture	64
4.3	Switch Interface	65
4.3.1	Input Models	65
4.3.2	Output Models	65
4.4	Cell Switch Fabric	65
4.5	Connection Admission Control (CAC)	66
4.6	Switch Management	66
4.7	The Cell Switch Fabric	67
4.7.1	Connection, Expansion And Multiplexing	68
4.7.2	Routing And Buffering	68
4.7.2.1	Shared Memory Approach	69
4.7.2.2	Shared Medium Approach	70
4.7.2.3	Fully Interconnected Approach	72
4.7.2.4	Space Division Approach	73
4.7.2.4.1	Banyan Networks	73
4.7.2.4.2	Delta Networks	75
4.7.2.4.2.1	Blocking And Buffering	75
4.7.2.4.2.2	Multiple-Path MINs	76

4.8	Switch Design Principles	77
4.8.1	Internal Blocking	77
4.8.2	Buffering Approaches	78
4.8.2.1	Input Queuing	79
4.8.2.2	Output Queuing	79
4.8.2.3	Internal Queuing	79
4.8.2.4	Recirculating Buffers	79
4.8.3	Buffer Sharing	80
4.8.4	Scalability Of Switch Fabrics	80
4.8.5	Multicasting	82
4.8.5.1	Shared Medium And Fully Interconnected Output-Buffered Approaches	83
4.8.5.2	Shared Memory Approach	83
4.8.5.3	Space division Fabrics	83
4.8.5.3.1	Crossbar Switches	83
4.8.5.3.2	Broadcast Banyan Networks	83
4.8.5.3.3	Copy Networks	84
4.8.6	Fault Tolerance	85
4.8.7	Using Priorities In Buffer Management	86
4.8.7.1	Preliminaries	86
4.8.7.2	Cell Scheduling	87
4.8.7.3	Cell Discarding	87
4.8.7.4	Congestion Indication	88
Chapter 5	ATM Technology Components: Upper Layers	89
5.1	ATM Adaptation Layer Functions	90
5.2	ATM Adaptation Layer Types	91
5.2.1	AAL Type 1 Services And Functions	91
5.2.2	AAL Type 2 Services And Functions	93
5.2.3	AAL Type 3 Services And Functions	95
5.2.4	AAL Type 4 Services And Functions	97
5.2.5	AAL Type 5 Services And Functions	97
5.3	ATM Services	99
5.3.1	Cell-Relay Services (CRS)	99

5.3.2	Frame-Relay Service	99
5.3.3	Connectionless Services (SMDS)	101
5.3.4	LAN Emulation Services (LES)	101
5.3.4.1	The Need For LAN Emulation	101
5.3.4.2	LAN-Specific characteristic To Be Emulated	102
5.3.4.2.1	Connectionless Services	102
5.3.4.2.2	Multicast Services	102
5.3.4.2.3	MAC Driver Interfaces In ATM Stations	102
5.3.4.2.4	Emulated LANs	103
5.3.4.2.5	Interconnection With Existing LANs	103
5.3.4.3	Description Of LAN Emulation Service	104
5.3.4.3.1	Basic Concepts	104
5.3.4.3.2	LAN Emulation In End Station	104
5.3.4.3.3	LAN Emulation In Bridges And LAN Switches	105
5.3.4.3.4	The Address Resolution Process	105
5.3.4.3.5	Forwarding Of Broadcast And Multicast Frames	106
5.3.4.3.6	The LAN Emulation Configuration Server	107
5.3.4.3.7	Network Architectures With LAN Emulation	108
Chapter 6	ATM And The Future	109
Conclusion		112
Appendix I		113
References		117

Acknowledgment

I would like to give a special thanks to Prof. Dr. Fakhreddin Mamedov the chairman of E.E. Engineering Department of Near East University, and my supervisor who cooperate with me and make it easy to do this project. For their vision of an institution for advanced scientific and engineering education and research, and for indefatigable efforts in realizing that vision in this young university.

Thanks to my family for their love and support all these years.

Also I would like to thank many people whose name may not appear on the cover, because without their cooperation, friendship, and understanding, producing this project would have been difficult.

M. A

ABSTRACT

Asynchronous Transfer Mode (ATM) is the optimal solution developed and accepted as the standard for public and private networks to implement and begin their evolution toward the Integrated Broadband Communication Network (IBCN). It is a flexible service available today.

Due to the direction of international standards it is important for every one to become familiar with the language and principles of ATM and begin to incorporate this information into their own planning process. This project is an in-depth introduction to what ATM is exactly, as well as where and why this technology originated and emerged as a revolutionary technological breakthrough. ATM is presented both in developmental and deliverable perspectives which are then contrasted with other services that are utilized today. This contrast clearly delineates the differences that exist between these current services and ATM, and high lights the real values that can be achieved with ATM as a major part of a user network. While some other types of telecommunication network Still exist at ATM is continuously taking place at the recent days. With its powerful and efficient aspects and attitudes and its considered as the living fact of the future of the world. We will discuss the aspects, characteristic functions, behaves, methods and designs of any given ATM network.

CHAPTER 1

Introduction TO Networking

Overview

First but not the last, This chapter will give an idea about networking, science and network expansion (from where it started and where it goes), the situation of networking in the past and present, networking issues and challenges going through its standardization to its failure and success. This chapter also talk a little bit about B-ISDN, that's explaining the meaning of B-ISDN and the idea of it why do we need it, its services and characteristic

1.1 Science and Network Expansion

Today there is an international computer network. It spans the globe and connects universities, researchers, computer workers and users around the world. Twenty-five years ago, these developments were non existent. This significant development has involved millions of people around the world. But others who are not participants in this exciting new global computer community know practically nothing of its existence.

The current network is the result of work done by scientists, engineers, programmers and other networking pioneers who functioned in the experimental tradition established by the Royal Society in London in the 1660s.

In a similar way, 300 years later, in 1969, work on the Arpanet began to give academic computer scientists and other U.S. Department of Defense contractors a way to scientifically test their networking theories. Based on an actual network to help them to collaborate and to test their theory and make it accord with the real problems of creating a computer network, a global network evolved which amazed even the pioneers themselves.

Community it made possible. Usenet News now reaches over 6 million people worldwide with over 4,500 different newsgroup subjects and gigabytes of articles. This news uses no paper, no glue, and no postage. It is created and distributed by a highly automated process. This technology makes it possible for the users themselves to determine and provide for the content and range of information conveyed via this new news medium. It also makes possible the rapid response and discussion of articles contributed by users and

provides a forum where issues can be freely debated and information exchanged. This news provides for the information exchange and learning needed by the system administrators, programmers, engineers, scientists and users to do their day to day work. In turn they contribute the programs and articles required for the network's development.

1.2 Issues and Challenges Ahead

1.2.1 Networking is Critical

Networking has become the most critical part of computing. Today, computers are used mostly for transferring information from one peripheral to another, from network to the disk, from disk to the video screen, from keyboard to the disk, and so on. Mail, file transfer, information browsing using World Wide Web, Gopher, and WAIS takes up more time of the computing resources than computing per se. Initially, when the computers were designed, the performance was measured by the "add" instruction time. Today, it is the "move" instruction that is the key to the perceived performance of a system. This means that the bus performance is more important than the arithmetic logical unit (ALU) performance. I/O performance is more important than the SPECmarks.

There are several other reasons for communications and networking becoming critical. First, the users have been moving away from the computer. In the sixties, computer users went to computer rooms to use them. In the seventies, they moved to terminal rooms away from the computer rooms. In the eighties, the users moved to their desktop. In the nineties, they are mobile and can be anywhere. This distance between the users and the computers has lead to a natural need for communication between the user and the computer or the user interface device (which may be a portable computer) and the servers.

Second, the system extent has been growing continuously. Up until eighties, the computers consisted of one node spread within 10 meters. In nineties, the systems consists of hundreds of nodes within a campus. The increasing extent leads to increasing needs for computing.

In the last ten years, we have seen increasing personalization of computing resources. We moved from timesharing to personal computing. Now we need ways to work

together with other users. So, in the next ten years, emphasis will be on cooperative computing. This will further lead to increase in communications.

In the last decade, we were busy developing corporate networks, and campus networks. In the next decade, we need to develop intercorporate networks, national information infrastructures, and international information infrastructures. All these developments will lead to more growth in the field of networking and more demands for the personnel with skills in networking.

The increasing role of communications in computing has lead to the merger of the telecommunications and computing industries. The line between voice and data communications is fading away. Data communication is expected to take over voice communication in terms of volume.

1.2.2 Peak of Technology Life Cycle.

Most technologies follow an S-shaped curve, where the number of problems solved is plotted against time. There are three distinct stages in the life of a technology. In the beginning, all problems are hard and it takes a lot of resources and time to solve a few problems. At this stage, a lot of money is spent in research but there is very little revenue. Most of this research is funded by the traditional government funding agencies, such as, National Science Foundation and Advanced Research Project Agency (ARPA) in the United States.

After some of the key problems have been solved, a lot of other problems can be solved by spending little money. At this stage, the curve takes an upturn. The amount of revenues to be made from the technology is much more than the investments. It is at this stage, that the industries take over technology development. Numerous small companies are formed and quickly grow to become large corporations.

Finally, when all the easy problems have been solved, the remaining problems are hard and would require a lot of resources. At this stage, the researchers usually move on to some other technology and a new S-shaped curve is born.

The computing industry in general and the networking sector in particular is currently going through the middle fast growing region of the technology life cycle curve. The number of problems solved is indicated by the deployment of the technology. In case

of networks, one can plot the number of hosts on the networks, bytes per host, number of networks on the internet, total capacity (in MIPS) of hosts on the network, total memory or total disk space of the hosts on the network. In each case, one would see a sudden exponential up turn in the last few years.

1.2.3 Standardization

When a technology reaches the middle fast growing region, it becomes necessary to standardize it to make it usable for the masses. The computing industry in general and networking in particular is undergoing through this phase. Even if people use different computers, it is necessary that the networking interfaces be standardized so that these diverse computers can communicate with one another.

The standardization requires a change in the way business is done. Before standardization, a majority of the market is vertical. The only way for users to maintain compatibility is to buy the complete system from one manufacturer. System vendors make more money than component vendors. IBM, DEC, and Sun Microsystems are examples of such system vendors. After standardization, the business situation changes. Users can and do buy components from different vendors. The market becomes horizontal. Companies specializing in specific components and fields take prominence. Intel for processors, Microsoft for operating systems, Novell for networking are examples of this trend.

To survive in this post-standardization era, invention alone is not sufficient. Only those new ideas that are backed by a number of vendors become standardized and are adopted. It, thus, becomes necessary to form technology partnerships.

1.2.4 Past Failures and Successes

In the last fifteen years, we have seen a number of networking technologies that were very promising during their lifetime but were not successful in the long run. A sample of such technologies is listed in Table 1. In each case, the technology listed in the second column was more promising than the one in the third column until the year shown in the first column.

In early 80's, when Ethernet was being introduced, some argued that broadband Ethernet, which allows voice, video, and data to share a single cable would be more

popular than baseband Ethernet. As we all know, today there are a few broadband installations. Most installations of Ethernet are baseband. The cost of combining the three services was just too high. The analog circuits required for frequency multiplexing were not as reliable and economical as digital circuits with separate wiring.

Around the same time, when computer companies were trying to sell Ethernets, PBX manufacturers were presenting PBX as the better alternative, again because it was already there and it could handle voice as well as data. However, PBX was not accepted by the customers simply because it did not provide enough bandwidth.

The Integrated Service Digital Network (ISDN) was standardized in 1984 and was very promising then. However its deployment has been much too slow. Even after ten years, it is not possible to get an ISDN connection at most places. Even at those places where it is available, the 64 kbps bandwidth provided by it is not sufficient for most data applications. For low bandwidth applications, modems on analog lines provide a better alternative. Modem technology has advanced much beyond expectation. Today, one can get 28.8 kbps and 56 kbps modems that work with all pervasive analog lines and do not have monthly charges associated with the extra ISDN line.

In 1986, IEEE 802.4 (token bus) was touted as a better alternative than IEEE 802.3/Ethernet for real time environment. It was said that Ethernet could not provide the delay guarantees required for manufacturing and industrial environments. Manufacturing Automation Protocol (MAP/TOP) was seen as the right solution. Today, IEEE 802.3/Ethernet is used in all such environments. Token buses are practically nonexistent. Up until 1988, ISO/OSI protocol stack was seen as the leading contender for networking everywhere. Networking researchers in most countries were implementing ISO/OSI protocols. The United States Government Open Systems Interconnection Profile (GOSIP) even made ISO/OSI a mandatory requirement for government purchases. Today, TCP/IP protocol stack dominates instead. The OSI protocols suffered from the common problems of standards: it had too many features. Any feature required by any application in the world needed to be supported by the standard. The protocols took too long to standardize and were quite complex. The "build before you standardize" philosophy of the TCP/IP protocol stack helped in its success.

Up until 1991, IEEE 802.6 Dual Queue Dual Bus (DQDB) was seen as a promising candidate for metropolitan area networks. It is no longer considered viable. The unfairness problem and general problems of bus architectures have made it undesirable.

Xpress Transfer Protocol (XTP) was designed as the high-performance alternative to TCP/IP. Protocol Engines -- the company leading the design of XTP declared bankruptcy in 1992.

Table 1: Networking Failures and Successes of the Past

Year	Failure	Success
1980	Broadband Ethernet	Baseband Ethernet
1981	PBX	Ethernet
1984	ISDN	Modems
1988	OSI	TCP/IP
1991	DQDB	
1992	XTP	TCP

1.2.4.1 Requirements for Success

There are several lessons to learn from the list presented in Table1. First, all technologies appear very promising when first proposed. However, not all survive. Those that survive meet all the following requirements.

1. Low cost
2. High performance
3. Killer application
4. Timely completion
5. Interoperability among various implementations of the same technology
6. Coexistence with existing (legacy) technology.

1.3 The Situation In The Telecommunication World

Today's telecommunication networks are characterized by specialization. This means that for every individual telecommunication service at least one network exists that transports this service. Each of these networks was specially designed for that specific service and is often not at all applicable to transporting another service. When designing the

network of the future, one must take into account all possible existing and future services. The networks of today are very specialized and suffer from a large number of disadvantages:

- Service Dependence

Each network is only capable of transporting one specific service.

- Inflexibility

Advances in audio, video and speech coding and compression algorithms and progress in VLSI technology influence the bit rate generated by a certain service and thus change the service requirements for the network. In the future new services with unknown requirements will appear. A specialized network has great difficulties in adapting to new services requirements.

- Inefficiency

The internal available resources are used inefficiently. Resources which are available in one network cannot be made available to other networks. It is very important that in the future only a single network will exist and that this network is service independent. This implies a single network capable of transporting all services, sharing all its available resources between the different services. It will have the following advantages:

- Flexible and future safe

A network capable of transporting all types of services that will be able to adapt itself to new needs. Efficient in the use of its available resources All available resources can be shared between all services, such that an optimal statistical sharing of the resources can be obtained.

- Less expensive

Since only one network needs to be designed, manufactured and maintained the overall costs of the design, manufacturing, operations and maintenance will be lower.

1.4 B-ISDN

The highest bit rate a 64 Kbit/s based ISDN can offer to the user about 1.5 Mbit/s or 2 Mbit /s, respectively, i.e. the H1 channel bit rate. Connection of local area networks (LANs),

however, or transmission of moving images with resolution may, in many cases, require considerably higher bit rates.

1.4.1 Definition of Broadband Services

The elements that are necessary to construct a broadband network are already available. The fundamental element is high capacity (bandwidth) fibre optic cable for use in both the trunk network and subscriber loops. In most instances these elements are being introduced for reasons other than to provide broadband services. To some extent, therefore, the implementation of an integrated broadband network is inevitable. Consequently it is necessary to identify the applications to be supported and services to be provided by such a network.

1.4.1.2 What Is Broadband?

Broadband is defined as the provision of subscriber access at bit rates in excess of 2 Mbit/s (or 1.5 Mbit/s in the United States). Broadband services are the facilities (switched or non-switched) that a network operator provides to support broadband applications via an integrated subscriber access, that is, a single network port that provides access to all services. Although today independent networks (e.g. video-conferencing networks provide some broadband services, integrated access will be an essential feature of the coming broadband network. Broadband applications are any end-uses of the broadband network capabilities. Typical applications include LAN interconnection, interactive video, video telephony and video- conferencing. In particular a network capability (service) may be used by many applications.

1.4.2 Characteristics of Broadband Services

As far as service classes are concerned services are first broadly divided into two major groups: interactive and distributive services:

- **Interactive services:** involve an exchange of information between the originating subscriber and the service object, which in some cases might be a machine, e.g. database, mail service, and not another subscriber;

- **Distributive services:** has an essentially passive "entertainment" quality. They may be further subdivided into conversational services, which involve the receiver and sender simultaneously (including point-to-point and multicast applications), retrieval services, which entail the extraction of stored information from a database, and messaging services, which involve both sending information to a database for later retrieval and user-to-user electronic communication.
- **Continuous bit oriented,** where continuous refers to the stream nature of the source, rather than indicating absolutely constant bit rates. Such sources require circuit-mode emulation and a connection-oriented protocol for call establishment.
- **Bursty bit oriented,** corresponding to more random sources that typically use a packet-like mode and require either a connection-oriented or connectionless protocol.

1.4.3 Broadband Services

Some of the concepts associated with broadband services can be clarified by considering the uses for these services and their requirements. The services considered here are significantly different or have no narrowband counterparts:

- **TV Distribution:** enables subscribers to receive one or more television programs of a chosen quality (e.g. high definition TV, standard TV) in real-time and without blocking. Normally access to the service will be in two phases: first, during the call establishment phase the subscriber selects a number of channels (which form a channel set) for viewing; second, the subscriber selects programs within this channel set using a simplified signaling procedure. TV distribution is targeted at entertainment and educational applications.
- **Hi-fi Distribution:** This service enables a subscriber to receive one or more high quality sound programs. Like TV distribution, it is targeted at both entertainment and educational applications.
- **Videotelephony:** is a real-time, unidirectional or bi-directional audio-visual telecommunication service that provides person-to-person communication for the transfer of voice, moving pictures, and scanned images between two locations. Videotelephony is envisaged in the broadest sense not only as a personal communication medium, but also

for the transfer of recorded video material, documents, etc. Potential applications are numerous and well documented. Typical uses are:

- Face-to-face individual or group communication
 - Still images (objects, pictures, etc.)
 - Instruction and education
 - Professional consultation
 - Playing games (interactive, multi-user games)
 - Business discussions (individuals, groups, meetings, conferences)
 - Interviews
 - Buying and selling
 - Remote monitoring
 - Video mail.
- **Video Retrieval Applications:** include video-on-demand, advertising, video publishing, video training, and similar usages. The common characteristic of all these applications is that video material is generally sent simultaneously to more than one subscriber. It is assumed that under most circumstances video material will not be transmitted in real-time, so it will have to be recorded at the subscriber's premises. This assumption is made in view of the enormous bandwidth requirements, in both the network and server, if real-time access is sought to the entire range of potential programs over a reasonable size service area. This is a clearly separate service from pay TV, which is a distributive service.
 - **Motion Videotext:** is the broadband equivalent of the existing videotext service. As such, video access to databases via the telecommunications network is foreseen using the same service components as for TV distribution. Potential applications include:
 - Education and training
 - Downloading software (telesoftware)
 - Teleshopping
 - News retrieval
 - Advertising.

Data Broadband: Data transfer requires conversational, messaging and retrieval services. The attributes and components of data-related services are as varied as the applications

themselves, although the data transfer unit length together with their required response time and link error performance is what typically distinguishes applications. B-ISDN is tailored to become the universal future network.

B-ISDN implementation will, according to the CCITT, be based on the asynchronous transfer mode (ATM). Which will be discussed in the next chapter. Let's just specify the establishment of it.

1.5 Service introduction

Demand for broadband services is in its initial phase and growing. While just at the beginning of the commercial offer by network operators and service providers, inherent speed, response time and quality advantages will ensure the gradual uptake of services based on broadband networks. Naturally services for which there is a latent demand will be the first used on a wide scale.

The principal requirement expressed by the business sector is for high-speed data interchange services. Immediate needs for the interconnection of LANs and terminals over local/ national, and (possibly) international distances are growing as many companies become more dispersed and at the same time increasingly dependent on shared data. Examples of this dependence are found in areas such as factory automation, paperless ordering and engineering manufacturing integration.

A secondary requirement comes from the increasing use of video-conferencing to replace business travel. The use of dedicated switched conference facilities is currently expanding and such facilities are certainly targets for the broadband switched network.

Given that the residential sector has not yet made much use of narrowband ISDN services, it is hard to see an immediate demand for broadband services. One exception however is the home entertainment sector. The demand for improved quality cable television and simultaneous growth in the installation of residential fibre access offers a chance for integration, at least at the access level, of residential communication and entertainment needs.

This will provide a starting point for future broadband services. Many studies based on questionnaires and analyses of cable TV services have shown that, in principle at least,

subscribers are also willing to pay for various video retrieval services and enhanced quality TV. Apart from the problems of stimulating demand for new broadband services, a number of other factors also affect service acceptance.

Many businesses are now operating private networks over national and international distances, generally at sub-broadband speeds. The nascent broadband networks must not only attract new subscribers for the services offered but must also convince present users of private networks to switch from the private to the public domain. Factors other than cost and speed must be addressed if this transfer is to take place, including:

- **Security:** The public network must offer secure data transfer, which might involve encryption services.
- **Convenience:** The new network must be convenient for private network operators to use, enabling them to administer their private applications with few restrictions using the public network as a resource.
- **Control:** Private network operators want some degree of control over their own applications (e.g. network configuration, time-of-day changes in routing), making it necessary for the public operator to offer a virtual private network solution.

At the same time, telecommunications equipment suppliers must address the need to add integrated broadband services to private networks.

1.6 ATM Forum

The ATM Forum was started in October of 1991 by a consortium of four computer and telecommunication vendors. Since its inception, it has seen unprecedented growth, and today (as of June 1994) has over 500 members. Today's membership is made up of network equipment providers, semiconductor manufacturers, service providers, carriers and, most recently, end users.

The Forum is not a Standards body. The ATM Forum is a consortium of companies that writes specifications to accelerate the definition of ATM technology. These specifications are then passed up to ITU-T (Formerly the CCITT) for approval. The ITU-T standard body fully recognizes the ATM Forum as a credible working group.

CHAPTER 2

Asynchronous Transfer Mode (ATM)

Overview

This chapter gives the explanation of ATM and its technical characteristic. It starts with explaining the general meaning of ATM, how does it work, who is the user of ATM, why do we need it and the benefits of it, also a brief idea about the problem addressed by ATM. It goes through the solution of ATM (The ATM cell), the structure of ATM cell, the cell behavior, size and formats.

2.1 What Is ATM?

ATM stands for Asynchronous Transfer Mode. It is considered as the ground on which B-ISDN is to be built (the transfer mode for implementing B-ISDN). Asynchronous Transfer Mode (ATM) is an extremely high speed, low delay, multiplexing and switching technology that can support any type of user traffic including voice, data, and video applications. ATM is ideally suited to applications that cannot tolerate time delay, as well as for transporting frame relay and IP traffic that are characterized as bursty.

Before ATM, separate networks were required to carry voice, data and video information. The unique profiles of these traffic types make significantly different demands on network speeds and resources. Data traffic can tolerate delay, but voice and video cannot. With ATM, however, all of these traffic types can be transmitted or transported across one network (from megabit to gigabit speeds), because ATM can adapt the transmission of cells to the information generated.

2.2 How Does ATM Work?

ATM works by breaking information into fixed length 53-byte data cells. The cells are transported over traditional wire or fiber optic networks at extremely high speeds. Because information can be moved in small, standard-sized cells, switching can take place

in the hardware of a network, which is much faster than software switching. As a result, cells are routed through the network in a predictable manner with very low delay and with determined time intervals between cells. For network users, this predictability means extremely fast, real-time transmission of information on the same network infrastructure that supports non-real-time traffic support.

ATM is a connection-oriented protocol; this means that ATM must establish a logical connection to a defined endpoint before this connection can transport data. Cells on each port are assigned a path and a channel identifier that indicates the path or channel over which the cell is to be routed. The connections are called virtual paths or virtual channels. These connections can be permanently established or they can be set up and released dynamically depending upon the requirements of the user. A path can be an end-to-end connection in itself, or it can be a logical association or bundle of virtual channels.

ATM creates a common way of transmitting any type of digital information from any other intelligent device over a system of networks. Many types of networks can be consolidated using one technology. This gives ATM a substantial advantage over other transport technologies.

2.3 Who Is Using ATM

ATM was designed for users and network providers who require guaranteed real-time transmission of voice, data, and images while also requiring efficient, high performance transport of bursty packet data. Hospitals are using ATM to share real-time video and images for long distance consultation during diagnosis and operations. Schools are using ATM to bring students and instructors together, regardless of their locations. Corporations whose employees are in different locations benefit by using ATM to effortlessly share even the largest data files. The Internet runs on high-speed ATM backbones.

ATM is also proving to be an important technology for managing the demands placed on overburdened networks by the surge of Internet traffic. By adding a rich layer of traffic engineering to the Internet protocol (IP), ATM supports the transportation of traffic according to priority level and class of service, neither of which are supported by IP alone.

The result is a more effective means of ensuring the transmission of mission-critical traffic, a growing concern of services providers and businesses alike.

2.4 Why Do We Need ATM ?

- **There are four basic reasons:**

1. ATM has grown out of the need for a worldwide standard to allow interoperability of information, regardless of the "end-system" or type of information. With ATM, the goal is one international standard. There is an unprecedented level of acceptance throughout the industry of both the technology and the standardization process. With ATM, we are seeing an "emerging technology" being driven by international consensus, not by a single vendor's view or strategy.
2. Historically, there have been separate methods used for the transmission of information among users on a Local Area Network (LAN), versus "users" on the Wide Area Network (WAN). This situation has added to the complexity of networking as user's needs for connectivity expand from the LAN to metropolitan, national, and finally worldwide connectivity. ATM is a method of communication which can be used as the basis for both LAN and WAN technologies. Over time, as ATM continues to be deployed, the line between local and wide networks will blur to form a seamless network based on one standard-ATM.
3. Today, in most instances, separate networks are used to carry voice, data and video information-mostly because these traffic types have different characteristics. For instance, data traffic tends to be "bursty"-not needing to communicate for an extended period of time and then needing to communicate large quantities of information as fast as possible. Voice and video, on the other hand, tend to be more even in the amount of information required-but are very sensitive to when and in what order the information arrives. With ATM, separate networks will not be required. ATM is the only standards based technology which has been designed from the beginning to accommodate the simultaneous transmission of data, voice and video.
4. As described above, ATM is the emerging standard for communications. This is possible because ATM is available at various speeds from Megabits to Gigabit speeds.

ATM was developed because of developing trends in the networking field. The most important parameter is the emergence of a large number of communication services with different, sometimes yet unknown requirements. In this information age, customers are requesting an ever increasing number of new services. The most famous communication services to appear in the future are HDTV (High Definition TV), video conferencing, high speed data transfer, videophony, video library, home education and video on demand. This large span of requirements introduces the need for one universal network, which is flexible enough to provide all of these services in the same way. Two other parameters are the fast evolution of the semi-conductor and optical technology and the evolution in system concept ideas-the shift of superfluous transport functions to the edge of the network.

2.4.1 ATM Benefits - I

One Network-ATM will provide a single network for all traffic types-voice, data, and video. ATM allows for the integration of networks improving efficiency and manageability. Enables new applications-Due to its high speed and the integration of traffic types, ATM will enable the creation and expansion of new applications such as multimedia to the desktop. Compatibility-Because ATM is not based on a specific type of physical transport, it is compatible with currently deployed physical networks. ATM can be transported over twisted pair, coax and fiber optics.

2.4.2 ATM Benefits - II

- Incremental Migration-Efforts within the standards organizations and the ATM Forum continue to assure that embedded networks will be able to gain the benefits of ATM incrementally-upgrading portions of the network based on new application requirements and business needs.
- Simplified Network Management-ATM is evolving into a standard technology for local, campus/backbone and public and private wide area services. This uniformity is intended to simplify network management by using the same technology for all levels of the network.

- Long Architectural Lifetime-The information systems and telecommunications industries are focusing and standardizing on ATM. ATM has been designed from the onset to be scalable and flexible in:
 - Geographic distance
 - Number of users
- Access and trunk bandwidths (As of today, the speeds range from Megabits to Gigabits)
- This flexibility and scalability assures that ATM will be around for a long time.

2.5 Problems Addressed By ATM

A transfer mode has been defined as previously as the main technique of transmitting, multiplexing, switching, and receiving information in a network. Every communication technology from the telephone to broadcast has had its own "transfer mode". The network transfer mode was specialized for network functions. This is the key to understanding what ATM is for. A network could not be built using voice packets, since this "transfer mode" was developed and optimized for data.

All these methods involve synchronous use of the network. ATM, in contrast, is structured to work in asynchronous manner. The terms synchronous and asynchronous, as applied to transfer mode, refer to the scheme of multiplexing: mixing traffic from many sources together on the same physical network path. In a synchronous transfer mode, each source is assigned a fixed bandwidth based on position: a frequency band in FDM or a time slot in TDM. ATM is not based on position in the data at all; a header identifies whose traffic it is and where it goes. All traffic is sent based on demand; no traffic, no bandwidth drain. Therefore, an ATM network is not service-dependent; it works well for voice and video as well as data. It is not inflexible; as bandwidth requirements for video decrease (the blue VCR "idle" screen, for example), ATM networks can easily adjust. It is not inefficient; resources assigned for now to a voice connection can be used later for data traffic. Everything in ATM is done based on connections, not channels, as is done in traditional time-division multiplexing.

2.5.1 Too Many Networks

One of the biggest problems with networks today is that there are too many of them. Corporations built SNA networks with the 1970s for their mainframe data needs and followed them with router-based networks for their LAN connectivity needs. The need to provide the company with low cost voice services led to the purchase of private branch exchanges (PBXs) and then to the development of a private tie-line (point-to-point voice channel) network to connect the PBXs. Video conferencing needs were addressed with still another network, and as more and more services become necessary for a corporation to remain competitive.

2.5.2 Too Many Services

All these various networks are based on different service parameters, such as bit rate and delay variation tolerance (figure 2.1). Connectionless services are those in which the receiver does not have to be connected in some way prior to communication.

TELEPHONE	DATA	CABLE TV	VIDEO CONF.
CONN-OR	CONN-LESS OR CONN-OR	CONN-LESS	CONN-OR
DELAY VAR SENSITIVE	DELAY VAR INSENSITIVE	DELAY VAR SENSITIVE	DELAY VAR SENSITIVE
LOW BW	LOW/HIGH BW	HIGH BW	LOW/HIGH BW
CBR	VBR	CBR	CBR & VBR

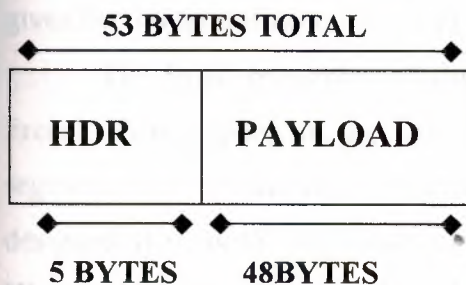
FIGURE 2.1

Information is merely packaged and sent as it is made available to the sender. Connection-oriented services require the establishment of a "connection" between sender and receiver. Delay variation sensitivity deals with the issue of how long data may be delayed within the network. Typically, this value will vary with the amount of traffic in the

network itself. When the network is busy, things are slower. The variation sensitivity of the receiver is an indication of how much this delay can vary (from 10 to 100 ms, for instance) before the receiver will think there is something wrong. Bandwidth just refers to the required bit rate of the service, which is highly dependent on whether some bit compression coding is used or not. Some services are always sending bits at a constant rate: constant-bit-rate (CBR) services. Some send bits in bursts: variable-bit-rate (VBR) services. Some networks offer the possibility of multiple quality of service (QOS) parameters that users may want at various times. The challenge of ATM is to make one physical network for all previous networks and services and be immune to change in service demands in the future.

2.6 ATM Solution: The ATM Cell

All ATM networking is based on the cell as the unit of data exchange (exchange). A cell is defined as a fixed-length block of information. Previous networks all used a simple stream of 0s and 1s that were organized into different network structures depending on the service and networks. This organization is still done with ATM networks, but at the endpoints of the networks. At the physical (bit) level, everything is sent and received as cell: a fixed-sized packet of bits.



***Payload may include some overhead bytes as well as data**

Figure 2.2 The ATM Cell

Although ATM networks are capable of integrating voice and video, the initial driving force will be the need for increased bandwidth and LAN interconnectivity. When it comes to the networking needs for data, rather than voice and video, ATM is designed to meet three main areas.

The first area is LAN interconnect bandwidth. As LANs run at higher and higher speeds and distributed computing becomes more common, existing digital speeds will become even more of a bottleneck. The second are of LAN networking efficiency in terms of the number of links needed to connect multiple LANs. Lastly, image and graphics applications pose their own problems due to the large sizes of these files that must cross the network.

2.7 The ATM Cell Structure

Asynchronous transfer mode (ATM) is easy to understand; it is simply a method of transferring information as it generated by a source using fixed-length cells. The asynchronous part refers to the "as it arrives" phrase in the definition. Cells are related to the concept of "cell relay."

2.7.1 The Cell

The basic idea of ATM is to segment data in small cells and then transfer them by the use of cell switching. Such cells have a uniform layout and a fixed size of 53 bytes, which greatly simplifies switching. Being more complex, packet-switching is not nearly fast enough to be of use for isochronous data (i.e. realtime video and sound). Cell-switching gives maximum utilization of the physical resources.

The basic principles of this technology were first formulated by AT&T and the French Telecompany in the first half of the 80's. Researchers found that if data was segmented in small, fixed size cells, switching could be done with simple, specially designed ICs. With sufficient intelligence to handle routing information in each cell, such ICs could be used to build very fast switch-matrixes. By connecting several switch-matrixes, highly efficient networks with small and predictable transmission delays can be built. The fact that ATM can be used efficiently in both WANs and LANs shows how powerful ATM is.

2.7.2 Connection-Oriented

When two participants wants to talk to each other, they first have to set up a connection between them. This connection remains open during the whole session, and

when the session is over the connection is closed. To the user it appears that the connection is permanent once established, but it's actually not. Cells are switched through the network so quickly that it appears as if there is a direct connection.

To be able to transmit a message to a host by ATM, a connection needs first to be made. ATM is thus connection-oriented.

2.7.3 Cell Size

An ATM cell has 5 bytes of header (administrative information) and 48 bytes of payload (data), counting a total of 53 bytes.

The choice of the cell size caused much controversy during CCITT standardization. The US computer branch wanted 64 bytes on payload, because they were considering the bandwidth utilization for data networks and efficient memory transfer (length of payload should be a power of 2 or at least a multiple of 4). 64 bytes fit both requirements.

Representants from the European and Japanese telephony branch was taking voice applications into consideration and wanted smaller cells. At cell sizes greater than 151, there is a talker echo problem. Cell sizes between 32 and 152 result in listener echo. Cell sizes less than or equal to 32 overcome both problems, under ideal conditions. Europeans telecommunication companies had no desire to invest money in echo canceling equipment, and thus went for 32 bytes of payload. However, because of big distances the US telecommunication companies had already installed echo cancellers across the country, and didn't consider this to be a problem.

After a lot of discussion 48 bytes of payload was agreed on. As far as the header goes, 10% of payload was perceived as an upper bound on the acceptable overhead, so 5 bytes was chosen. Although the header of 5 bytes is much for a 53 byte cell, the flexibility is worth the cost.

2.7.4 Cell Format

The header of 5 bytes contains 24 bits of VCI label, 8 bits of control field, and 8 bits of checksum. A possible structure of the ATM cell header is shown in the **Fig.2.3**.

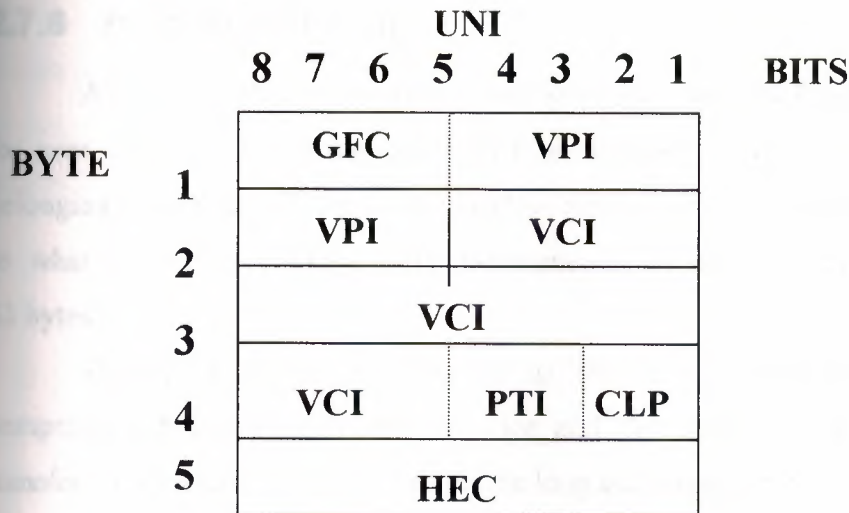


Figure.2.3 ATM Cell Header

GFC: GENERIC FLOW CONTROL.**PTI:**PAYLOAD TYPE INDECATOR**VPI:** VIRTUAL PATH IDENTIFIER.**CLP:**CELL LOSS PRIORITY**VCI:** VIRTUAL CHANNEL Identifier

ATM is sometimes referred to as label multiplexing or even asynchronous time-division multiplexing in older documentation. Both terms were used to refer to ATM when the standards were still under development. The label is the connection identifier that tells the receiver which connection the cell is to be associated with.

2.7.5 Time-Division Multiplexing

TDM works by having a fixed-length time slot assigned to each user input. This is therefore synchronous transfer mode (STM), since each slot is synchronized to a user input time slot. There is no need to identify the user bits in the data stream. IN STM, the ownership of bits is determined by position in the data stream. If no user bits have arrived to fill the time slot, it can not be given to another user. A special " idle bit pattern must be sent to each channel to keep the sender and receiver synchronized. This method can result in a great deal of idle bandwidth.

2.7.6 ATM Multiplexing

ATM multiplexing works by having a fixed number of cells per unit time available for user data. Each cell has the 5-byte header whose primary purpose is to identify cells belonging to the same "virtual channel" or connection. The cells are transmitted according to what is called in ATM the user's instantaneous real need. The cell length in ATM is only 53 bytes.

The cell length in ATM is set so small for a number of reasons. It is basically a compromise between the needs of voice and the needs of data applications such as file transfer. The whole idea is to avoid the long and unpredictable delays in waiting for long packets to finish transmission. ATM is aimed primarily at networks built on links running at 155 Mbps, Known as STS-3c speed, a 53-byte cell lasts only about 2.7 μ s: $(53 \text{ bytes} \times 8 \text{ bits/B}) / 155.52 \text{ MBPS} = 2.726 \mu\text{s}$.

2.7.7 ATM bandwidth and cell allocation

The aggregate bandwidth of all the ATM switches at the edges of the network is typically greater than the bandwidth available across the network. However, if the bandwidth of all the cell streams arriving at an intermediate ATM node exceeds its available outgoing band-width, that node is pen-nitted to discard any empty ATM cells, and use the outgoing band-width that would have been consumed by those empty cells to send full cells from some other cell stream.

This ability to discard empty cells, allows a network to support a large number of bursty cell sources. The empty cells are simply discarded wherever additional bandwidth is needed within the network.

The ability to insert or discard empty cells also makes it possible to join ATM carrier circuits operating at significantly different rates, using an ATM switch or multiplexor. This ability to easily adapt to varying rates allows ATM to scale from relatively low rates up to extremely high rates.

The main limiting factor at both ends becomes cost. Below T1 speeds, (1.554 Mbits/sec) ATM's high overhead ratio makes it progressively less cost effective than protocols speci-fically intended for lower-bandwidth operation. At extremely high speeds,

the cost and complexity of building very large and very fast switching hardware becomes a limiting factor. However, it is possible to use multiple parallel trunks and switches, to provide effective transmission rates much higher than the maximum rate for a single trunk and switch.

ATM System Architecture

ATM is a packet-switched network technology that is designed to support a wide range of services, including voice, video, and data. It is a connection-oriented technology, meaning that a connection must be established between the sender and the receiver before data can be transmitted.

ATM is a packet-switched network technology that is designed to support a wide range of services, including voice, video, and data. It is a connection-oriented technology, meaning that a connection must be established between the sender and the receiver before data can be transmitted. The ATM network architecture is based on the concept of a virtual circuit. A virtual circuit is a logical connection between two nodes in the network. It is established by the network controller and is used to route data packets from the sender to the receiver. The ATM network architecture is based on the concept of a virtual circuit. A virtual circuit is a logical connection between two nodes in the network. It is established by the network controller and is used to route data packets from the sender to the receiver.

3.1 ATM Network Architecture

The ATM network architecture is based on the concept of a virtual circuit. A virtual circuit is a logical connection between two nodes in the network. It is established by the network controller and is used to route data packets from the sender to the receiver. The ATM network architecture is based on the concept of a virtual circuit. A virtual circuit is a logical connection between two nodes in the network. It is established by the network controller and is used to route data packets from the sender to the receiver.

CHAPTER 3

ATM NETWORK CONCEPTS AND ARCHITECTURE

Overview

ATM is a new network architecture, but it is still very much a complete architecture. This chapter takes a look at the over all structure of ATM network, it will introduce the layers and detail their functions, it will look at how ATM networks must still address the issues of signaling, network performance, and traffic control, just as all network architectures have had to in the past. ATM will require many more capabilities in the area. ATM adds the requirement of addressing functions such as congestion control performance and services. We will also talk about ATM layers protocols and concepts.

ATM System Architecture

ATM is a layered architecture allowing multiple services like voice, data and video, to be mixed over the network.

As we explained in chapter 2, ATM is founded on the concept of a cell as the unit of information transfer. A cell is defined as a small fixed-length block, as opposed to the variable-length packet most data services use. The fixed length makes the cell useful for transferring other services such as voice and video at the same time over the same links. In other word we can define a cell as " a fixed-length block for supporting multiple quality-of- service parameters to different users over the same unchannelized physical network. ATM is still essentially a connection-oriented system, like the phone network, although it is intended that data services that are traditionally connectionless, such as LAN-to-LAN inter-connectivity, will be easy to use ATM as a transport easily.

3.1 ATM's Position In The OSI Reference Model

The open systems interconnect reference model (OSI-RM) was established by the international standards organization (ISO) in 1979. Seven layer models for communications over a wide area network (WAN), it was soon modified for functioning over a local area

network (LAN) as well. The problem was that the original data link layer (layer2) specification worked only between adjacent systems. Systems connected directly to each other over a single hop, point-to-point link the OSI-RM data link layer only needed a very few addresses for the frames, or protocol data units (PDUs), at layer2. The invention of LANs meant that there were no longer a small number of adjacent systems. Rather, every system on a LAN was adjacent to all the other systems due to the shared media nature inherent in LANs. This meant that with a LAN such as Ethernet, with up to 1024 adjacent systems, the available layer2 addresses in the original WAN OSI-RM were quickly exhausted long before any thing like 1024 systems could be attached. The model had to be modified to provide a solution. This modification consisted of dividing the data link layer (layer2) into an upper and lower portion: layer2b became the logical link control (LLC) sublayer, and layer2a became the media access control (MAC) sublayer.

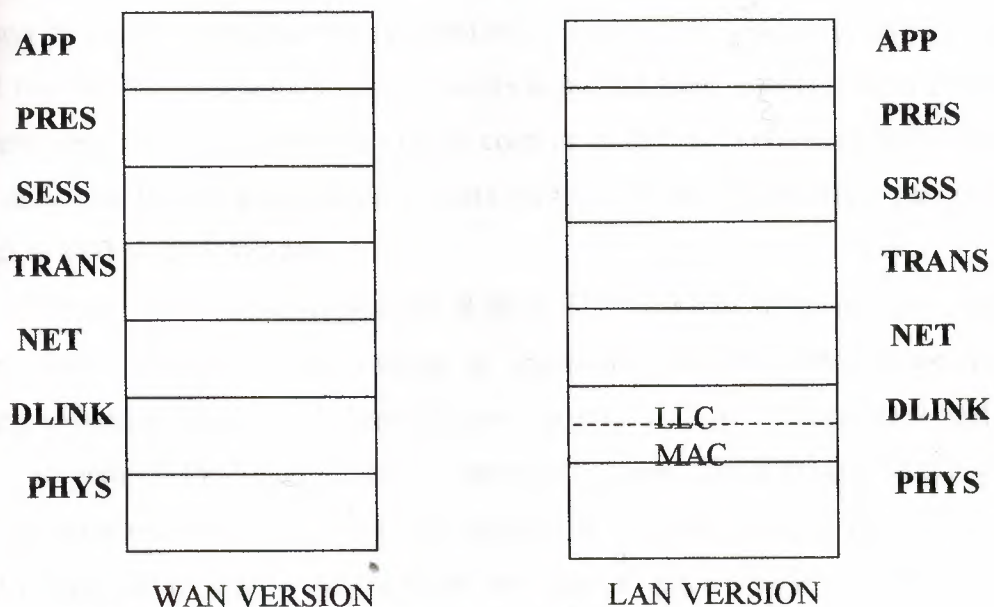


Figure 3.1 The OSI Reference Model (OSI-RM)

This splitting of layers is not uncommon and can be extremely useful in the implementation of these protocols in hardware and software. Layer2b, the logical link control, retained the function of the original WAN data link layer, which now became the "logical" link, since there were no longer any physical adjacent links needed for LANs.

Under2b, layer2a (the MAC sublayer) generated the proper frame structure and protocol for the various LAN technologies being developed in the early 1980s.

3.1.1 Router Based Networking

The most common paradigm for networking today is the concept of router based networking. All networks are connected by routers rather than other network connectivity devices such as bridges or gateways. Bridges or gateways are used for connectivity in many situations, bridges take a packet from one LAN, transport it verbatim across the WAN, and then transmit it over the remote LAN, just as if it had originated from some computer on the destination LAN. This approach provides only a limited solution, because it is not scaleable. Each connection between LANs requires a bridge at each end, and a dedicated point-to-point WAN link between those bridges. Additionally, the network layer protocols of the first LAN must match the network layer protocols of the second LAN. This problem was overcome by increasing the capabilities of bridges by combining them with devices called routers. Routers can connect networks using the same transport layer protocols, but different network layer protocols. These combined devices, known as bridge/routers (or more simply as brouters) can route packets between multiple LANs over packet switched, as well as TDM based, WANS.

Routers have a unique position in the OSI-RM. Many different types of networks maybe connected, with the routing of messages between them done by special relaying/switching devices. If these devices operate at layer3 of the OSI-RM and use layer3c to perform the routing then this device is a router. These devices handle layer three protocol data units (PDUs) known as packets if the network services is connection-oriented and datagrams if the network services is connectionless. These services are distinguished by the need in connection-oriented networks for a signaling protocol to establish the connection in the network before data can flow from a sender to a receiver. No signaling is needed for a data gram services it is common today to refer to layer3 network devices that handle packets as switches and to call layer3 network devices that handle datagrams as routers.

3.1.2 X.25 Packet-Switched Network

The term packets usually indicate a connection-oriented network service. The older (1984) X.25 standard describe a network architecture that looks very much like router-based networking. In X.25, an X.21 connector specifies the physical layer link access protocol- balanced (LAP-B), a part of high level data link control (HDLC, very similar to SDLC) forms the data link layer. The X.25 packet layer protocol (PLP) is the network layer the switching function in the packet switch device is not really a layer at all. It is really the over all function of the switch and is included merely to provide a correspondence to the routing function of a router.

3.1.3 Switching Versus Routing

The differences between router-based networking and switched-based networking are more in terms of operation than in architecture. Both network nodes operates at the bottom three layers of the OSI-RM; both involve taking in a layer3 PDU form an input port, looking it up in a table, and forwarding it through the network to the next network node, and so on.

The differences are not in the architecture but in the fact that routers are connectionless and switches are connection-oriented. Routers rout (connectionless layer3 PDUs). All layer3 PDUs have similar structure. They differ only in the specifics of field and lengths. A packet is a connection-oriented datagram, and a datagram is a connectionless packet. And by extension, a switch is a connection-oriented router, and a router is a connectionless switch.

Routers and switches do differ significantly in operation. Switches setup paths through the network and connection time. The connection time maybe done by contract at service provision time or upon processing the first datagram for connectionless services. Both switches and router may deliver connectionless and/ or connection-oriented services, connection-oriented switches will always be able to offer connection-oriented services in a simpler, more efficient, and more economical fashion. This path also may be set up using a special protocol designed just for the purpose: the signaling protocol.

Once the connection is setup, all traffic follows the same path through the network. There is only a simple table look up needed to switch traffic to the correct output. The table entries are established and maintained signaling protocol. Because all traffic follows the same path, there can be no out-of-sequence delivery. This has an important effect. It means that if a destination receive packet1 and then packet3 from the source and at the other end of the connection, the receiver knows that packet2 is not coming and can take immediate steps to correct the situation. A lost path may mean that all connection using the path maybe lost. And even if the switches can some how move the connection to a new path, the signaling protocol must be fast enough to update all the affected tables without losing much data.

Switches are much faster internally than routers. A router can work in a totally connectionless manner this requires a full routing algorithm or set of rules, to follow at each network node. Router tables are more complex than switch tables because they must take these possibilities into account.

Unlike switches these routers tables are maintained by a full routing protocol that runs between each adjacent routers in the network. Signaling protocols tend to be very rudimentary compared with routing protocols but routers gain the ability to dynamically rout traffic around failed links. Routers deliver datagrams out of sequence all the time since the actual rout a datagram takes may vary from that a receiver that has gotten datagram1 and datagram3 can make no assumption at all about datagram2. It may still be on its way by a longer rout or may never arrive. Delays are added at the destination to deal with these situations. And, of course, as soon as the destination notifies the sender to resent datagram2, the missing datagram2 shows up. Many network devices today combine aspects of switches and routers as designers seek to take advantage of the pluses of both.

3.1.4 Frame-Relay Network

Frame relay network takes the function of X.25, which include full error control and flow control hop by hop-- between switches-through the network and strip them down to the beer essentials-routing. Error control and flow control moves to the " edges" of the network. Frame relay is therefor a version of X.25 for the 1990s. Instead of the normal nine

processing steps that an X.25 switch must take to move a packet through a switch or network node, frame relay network nodes take only two processing steps. This makes the nodal processing delay much less than in X.25 network, making frame-relay much more suitable for the high speed network needed today.

Frame-relay networks have with both routers and X.25 networks; the NNI is very permissive, and routing-switching must still be done some how in the network. In frame-relay a logical connection, an identifier included in the frame but it is a connection nonetheless. This is the data link connection identifier (DLCI) in frame relay, a composite field made up of two separate fields from the original high layer data link control (DLC) frame structure that forms the basis for frame relay.

3.2 B-ISDN Protocol Reference Model

As stated earlier, ATM was originally developed as the switching structure for B-ISDN. Since that time ATM has found many other applications. This section will briefly discuss how the ATM switch fits into the B-ISDN network. The B-ISDN protocol reference model defined by ITU-T is shown in figure.

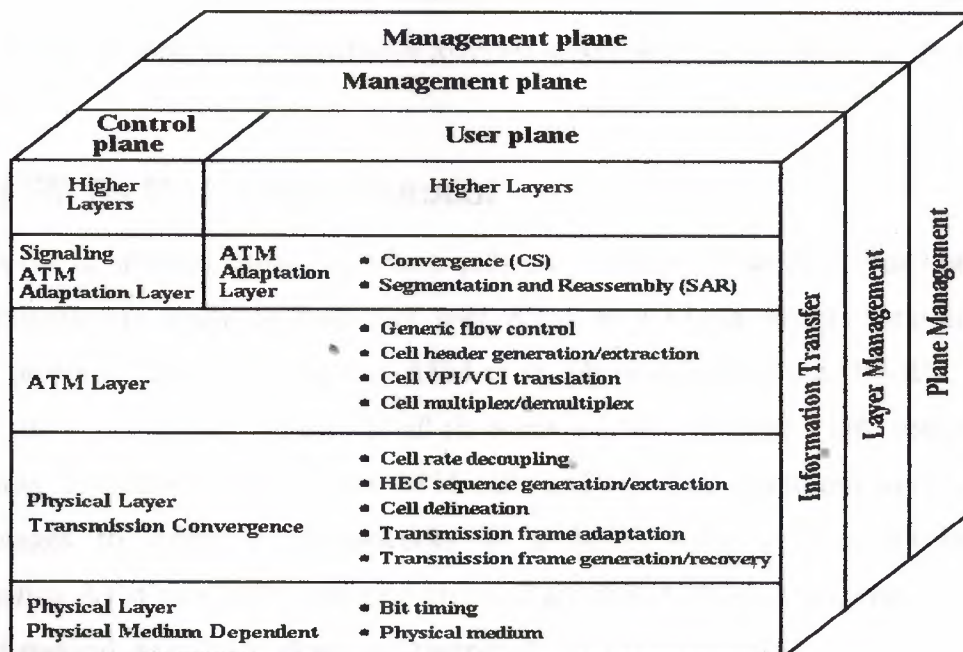


Figure3.2: B-ISDN Protocol Reference Model

it is composed of a user plane, a control plane and a management plane, each of them responsible for associated functions. In particular:

The user plane

With its layered structure, provides for user information flow transfer, along with associated controls ranging from flow control to error recovery, etc.;

The control plane

Has a layered structure and performs the call control and connection control functions; it deals with the signaling necessary to set up, supervise and release calls and connections.

The management plane

Provides two different types of functions:

"plane management" functions, not layered, that are related to a system as a whole and provide co-ordination between all the planes and "layer management" functions, that are related to resources and parameters residing in its protocol entities; layer management handles the operation and maintenance (OAM) information flows specific to the layer concerned.

3.2.1 ATM As MAC Layer Protocol

ATM capabilities depending where cells are available for services. For instance, an ATM network can easily be built that uses ATM as a Media Access Control (MAC) sublayer protocol. That is, it uses only ATM as the physical layer of the OSI-RM. There is no need for a LAN MAC sublayer at all (it is not a LAN), so logical link control (LLC) frames may be mapped directly into ATM cells and sent out. There are advantages and disadvantages to using ATM networks as a MAC sublayer. The advantages of implementing ATM as a MAC sublayer protocol are that ATM now becomes just another transport method, like Token Ring or Ethernet. It extends the OSI-RM. Mainly by adding the ATM sublayer to the physical layer and the ATM network access is totally transparent to the user. Where the disadvantages however, there is obviously no broadcast capability,

since ATM is connection-oriented. There is no access to QOS parameters by a user, since these are employed at higher layers in the OSI-RM.

3.2.2 ATM As a Link Layer Protocol

Network layer packets (IP data gram) maybe loaded into cells and sent across the ATM networks. There is no need for any other data link layer protocols at all. The LLC and MAC sublayers are not needed. There are advantages and disadvantages here as well. With two of the advantages listed before, ATM offering transparent user access and expanding the OSI-RM. But ATM is no longer just another transport method because of the way the network layer interacts with the network layer. That is, the unique network address exists at the network layer, but frames are sent based on the link layer address. Link layer PDUs must be sent as link layer PDUs based on the link layer address.

The disadvantage is that the network layer protocol has to be modified. ATM is a new way of doing networking, not just different. IP and other older protocols will have to be changed a little to work with ATM as a data link protocol.

3.2.3 ATM As a Network Layer Protocol

At the network layer, ATM can interface directly with the end-to-end layer: the transport layer. ATM forms the entire transport network from one endpoint to the other. The transport layer "address" is used, not the network layer address. ATM forming the entire transport network seems completely natural, and the popular convention of defining a cell as a fixed-length packet seems to identify ATM cells with layer 3 (network layer) PDUs (packets). Only one universal network address is needed, and only one routing

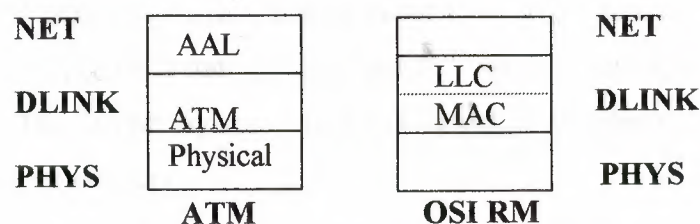


Figure 3.3.1 ATM as a network layer protocol

method is needed: ATM routing. Here the transport layer must be modified for an ATM interface.

3.2.4 ATM As Transport Layer Protocol

Here, with all applications are ATM applications now. Programs send and receive cells directly. There will be one application program Interface (API) for all applications: an ATM API. But this has disadvantages, such as direct application access to ATM and QOS access for multimedia and video, and it does make the most effective and efficient use of ATM networks. This requires major changes to existing program APIs and methodologies.

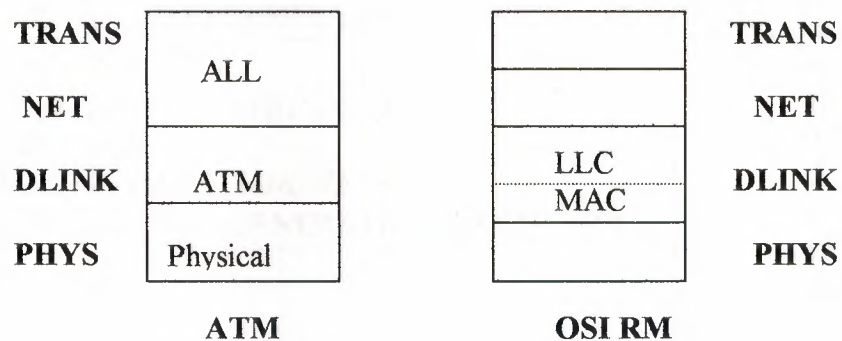


Figure3.3.2 ATM as a transport layer protocol

3.3 ATM Functions and Layers

The layers of the ATM protocol stack and the major functions performed at these layers are shown in Fig.3.4. The layer management is a function that spans all the layers and the use of various convergence components in several layers.

Convergence is an important ATM concept. It means that there are multiple options that may be employed above or below some layers in the model. Bits maybe framed or sent "raw" or sent on fiber or coaxial cable or come from any way else, however, the convergence layers help to present a uniform interface to other adjacent layers to make it easier for implementation. The lowest layer of the ATM model is the physical layer, Physical Layer is divided into two sub-layers.

L A Y E R M A N A G E M E N T	HIGHER LAYER FUNCTIONS	HIGHER LAYER	
	CONVERGENCE	C S	A A
	SEGMENTATION AND RESEMBLY	S A R	L
	GENERIC FLOW CONTROL CELL HEADER GENERATION/EXTRACTION CELL VPI/VCI TRANSLATION CELL MULTIPLEX AND DEMULTIPLEX	A T M	
	CELL RATE DECOUPLING HEC SEQUENCE GENERATION/VERIFICATION CELL DELINATION TRANSMISSION FRAME ADAPTATION TRANSMISSION FRAME GENERATION/ RECOVERY	T C	P H Y S I C A L
	BIT TIMING PHYSICAL MEDIUM	P M	L A Y E R

Figure.3.4. ATM function and layers.

1. Physical Medium (PM) sublayer**2. Transmission Convergence (TC) sublayer.**

The PM sublayer contains only the Physical Medium dependent functions. It provides bit transmission capability including bit alignment. It performs Line coding and also electrical/ optical conversion if necessary. Optical fiber will be the physical medium and in some cases, coaxial and twisted pair cables are also used. It includes **bit timing** functions such as the generation and reception of waveforms suitable for the medium and also insertion and extraction of bit timing information.

The **TC** sublayer mainly does five functions as shown in the figure. The lowest function is generation and recovery of the transmission frame.

The next function i.e. **transmission frame adaptation** takes care of all actions to adapt the cell flow according to the used payload structure of the transmission system in the sending direction. It extracts the cell flow from the transmission frame in the receiving direction. The frame can be a synchronous digital hierarchy (SDH) envelope or an envelope according to ITU-T Recommendation G.703.

Cell delineation function enables the receiver to recover the cell boundaries. Scrambling and Descrambling are to be done in the information field of a cell before the transmission and reception respectively to protect the cell delineation mechanism.

The **HEC** sequence generation / verification is done in the transmit direction and its value is recalculated and compared with the received value and thus used in correcting the header errors. If the header errors can not be corrected, the cell will be discarded.

Cell rate decoupling inserts the idle cells in the transmitting direction in order to adapt the rate of the ATM cells to the payload capacity of the transmission system. It suppresses all idle cells in the receiving direction. Only assigned and unassigned cells are passed to the ATM layer.

3.4 ATM Layer Functions

ATM layer is the layer above the physical layer. As shown in the figure, it does the 4 functions which can be explained as follows.

Cell header generation/extraction: This function adds the appropriate ATM cell header (except for the HEC value) to the received cell information field from the AAL in the transmit direction. VPI/VCI values are obtained by translation from the SAP identifier. It does opposite i.e. removes cell header in the receive direction. Only cell information field is passed to the AAL.

Cell multiplex and demultiplex: This function multiplexes cells from individual VPs and VCs into one resulting cell stream in the transmit direction. It divides the arriving cell stream into individual cell flows with respect to VC or VP in the receive direction.

VPI and VCI translation: This function is performed at the ATM switching and/or cross-connect nodes. At the VP switch, the value of the VPI field of each incoming cell is translated into a new VPI value of the outgoing cell. The values of VPI and VCI are translated into new values at a VC switch.

Generic Flow Control (GFC): This function supports control of the ATM traffic flow in a customer network. This is defined at the B-ISDN User-to-network interface (UNI).

3.5 ATM Adaptation Layer Functions (AAL):

AAL is divided into two sub-layers as shown in the figure 3.4.

1. Segmentation and Reassembly (SAR)

2. Convergence sublayer (CS)

1. SAR sublayer: This layer performs segmentation of the higher layer information into a size suitable for the payload of the ATM cells of a virtual connection and at the receive side, it reassembles the contents of the cells of a virtual connection into data units to be delivered to the higher layers.

2. CS sublayer: This layer performs functions like message identification and time/clock recovery. This layer is further divided into Common part convergence sublayer (CPCS) and a Service specific convergence sublayer (SSCS) to support data transport over ATM. AAL service data units are transported from one AAL service access point (SAP) to one or more others through the ATM network. The AAL users can select a given AAL-SAP associated with the QOS required to transport the AAL-SDU. There are 5 AALs have been defined, one for each class of service. The service classification for AAL can be shown by the following figure.

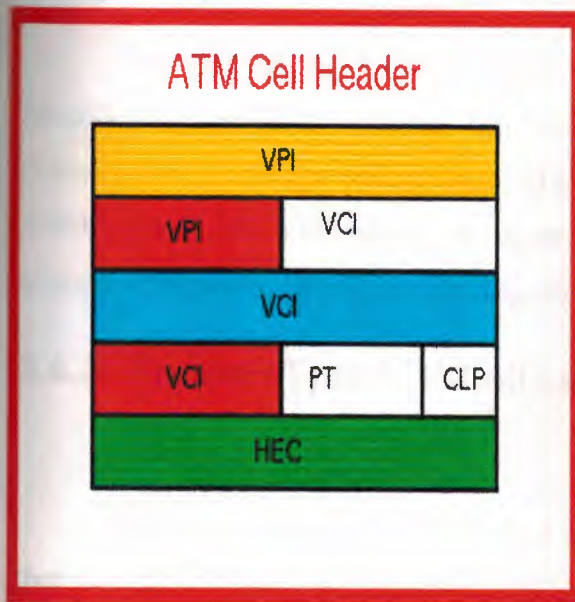
	CLASS A	CLASS B	CLASS C	CLASS D
Timing Relation Between Source and Destination.	Required		Not Required	
Bit Rate	Constant	Variable		
Connection Mode	Connection Oriented			Conne- ction Less

Figure.3.5 ClassX : Unrestricted (bit rate variable, connection oriented or Connectionless).

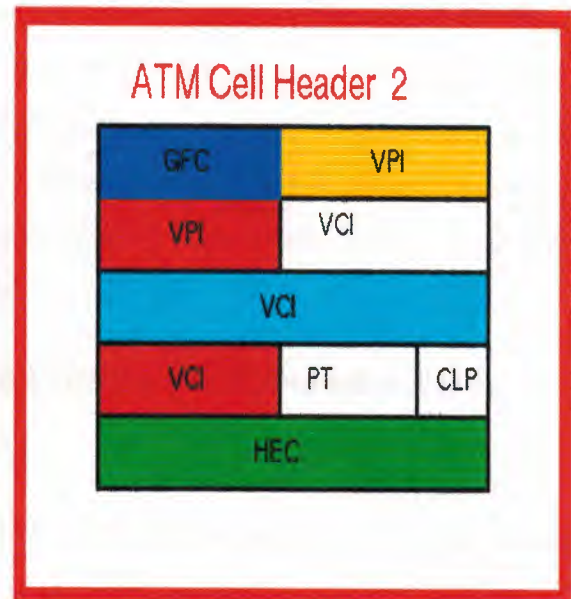
3.6 ATM Cell Structure Details

The ATM cell is the basic unit of information transfer in the B-ISDN ATM protocol. The cell is comprised of 53 bytes. Five of the bytes make up the header field and the remaining 48 bytes form the user information field. The following is the structure of the Network Node Interface (NNI) ATM and (UNI) ATM Cell Headers:

(NNI) ATM Cell Headers



(UNI) ATM Cell Headers



The header field is divided into GFC, VPI and VCI, PT, CLP and HEC fields. The associated bit sizes differ slightly at the NNI and the UNI. The bit sizes are as follows.

Bit Allocation of Cell Header

Function	UNI	NNI
GFC	4	0
VPI	8	12
VCI	16	16
PT	3	3
CLP	1	1
HEC	8	8

3.6.1 Generic Flow Control (GFC):

Although the primary function of this header is the physical access control, it is often used to reduce cell jitters in CBR services, assign fair capacity for VBR services, and to control traffic for VBR flows. Such functionality requires the power to control any UNI structure, be it a ring, a star, some bus configuration, or any combination of these.

3.6.2 Virtual Path Identifier / Virtual Channel Identifier (VPI/VCI):

The role of the VPI/VCI fields is to indicate Virtual path or virtual channel identification numbers, so that the cells belonging to the same connection can be distinguished. A unique and separate VPI/VCI identifier is assigned in advance to indicate which type of cell is following, unassigned cells, physical layer OAM cells, metasignalling channel or a generic broadcast signaling channel.

3.6.3 Payload Type (PT)/Cell Loss Priority (CLP)/Header Error

Control (HEC):

When user information is present or the ATM cell has suffered traffic congestion then the PT field will yield this information.

The CLP bit is used to tell the system whether the corresponding byte is to be discarded during network congestion. ATM cells with CLP=0 have a priority in regard to cell loss than ATM cells with CLP=1. Therefore, during resource congestions, CLP=1 cells are dropped before any CLP=0 cell is dropped.

HEC is a CRC byte for the cell header field and is used for sensing and correcting cell errors and in delineating the cell header.

The ATM cells traveling through the virtual pipes and channels mentioned above will need a physical medium (e.g. fiber optic cable) to go anywhere. So another layer in the B-ISDN protocol reference model was created, below the ATM layer (see the diagram above).

3.7 Virtual Channels/Paths

ATM provides two types of transport connection, Virtual paths and Virtual channels. A virtual channel is a unidirectional pipe made up from the concatenation of a sequence of connection elements. A virtual path consists of a set of these channels (Fig.3.6 below).

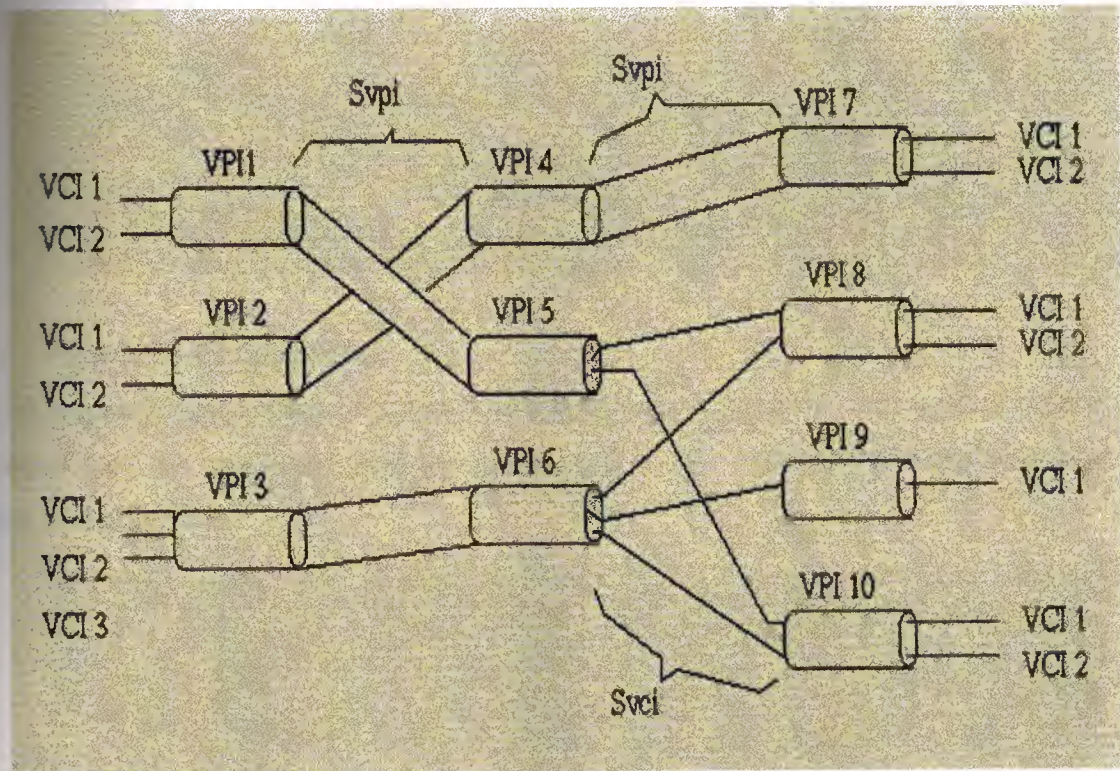


Figure3.6.1 a set of virtual path channels

Pa Each channel and path has an identifier associated with it. All channels within a single path must have distinct channel identifiers but may have the same channel identifier as channels in different virtual paths. An individual channel can therefore be uniquely identified by its virtual channel and virtual path number.

The virtual channel and path numbers of a connection may differ from source to destination if the connection is switched at some point within the network. Virtual channels which remain within the single virtual path through out the connection will have identical virtual channels identifiers at both ends.

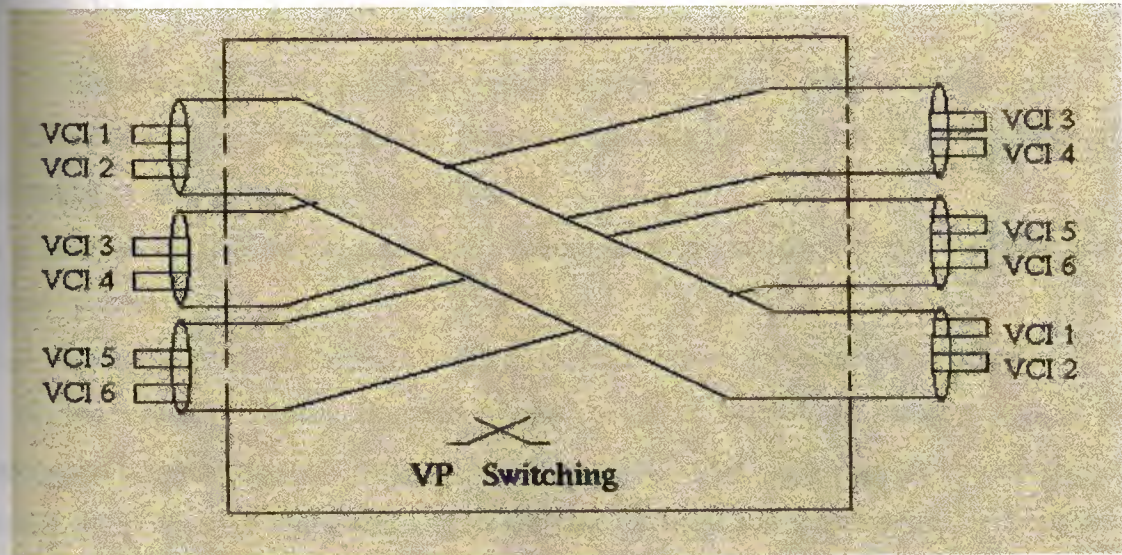


Figure 3.6.2 VP Switching

Cell sequence is maintained through a virtual channel connection. Each virtual channel and virtual path has negotiated QOS associated with it. This parameter includes values for cell loss and cell delay.

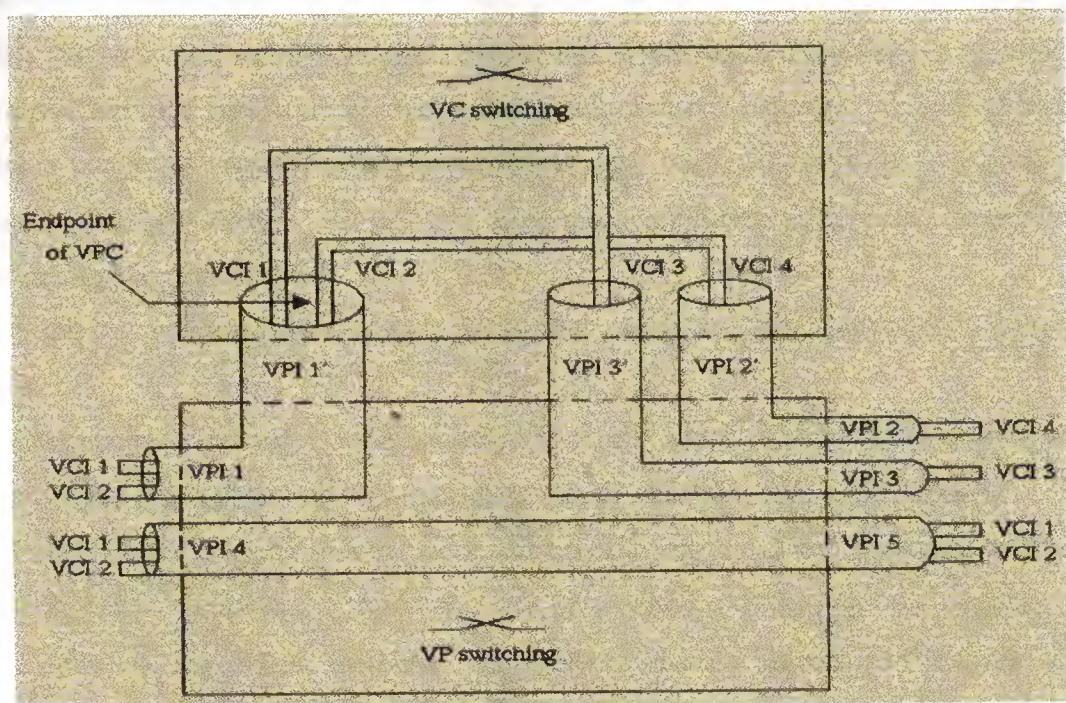


Figure 3.6.3 VP and VC switching

3.7.1 Virtual path/channel connection:

There are four ways in which a virtual channel or virtual path may be set up.

- 1) The virtual channel/path maybe pre-reserved with the network as in the case of a permanent or semi-permanent connection.
 - 2) A new connection maybe setup via a metasignaling procedure across a metasignaling virtual channel.
 - 3) A connection maybe set up as a result of a user to network signaling procedure.
 - 4) A new virtual channel connection may be setup within an existing virtual path connection between two user network interfaces using a user to user signaling procedure.
- During setup the user negotiates a QOS with the network and traffic parameters are setup. The network monitors the traffic to ensure this agreement is adhered to.

3.8 ATM Congestion Control

3.8.1 Why The Need Of Congestion Control

The assumption that statistical multiplexing can be used to improve the link utilization is that the users do not take their peak rate values simultaneously. But since the traffic demands are stochastic and cannot be predicted, congestion is unavoidable. Whenever the total input rate is greater than the output link capacity, congestion happens. Under a congestion situation, the queue length may become very large in a short time, resulting in buffer overflow and cell loss. So congestion control is necessary to ensure that users get the negotiated QOS.

There are several misunderstandings about the cause and the solutions of congestion control.

1. Congestion is caused by the shortage of buffer space. The problem will be solved when the cost of memory becomes cheap enough to allow very large memory. Larger buffer is useful only for very short-term congestions and will cause undesirable long delays. Suppose the total input rate of a switch is 1Mbps and the capacity of the output link is 0.5Mbps, the buffer will overflow after 16 second with 1Mbyte memory and will also overflow after 1 hour with 225Mbyte memory if the situation persists. Thus larger buffer size can only postpone the discarding of cells but cannot prevent it.

The long queue and long delay introduced by large memory is undesirable for some applications.

2. Congestion is caused by slow links. The problem will be solved when high-speed links become available.

It is not always the case; sometimes increases in link bandwidth can aggravate the congestion problem because higher speed links may make the network more unbalanced. For the configuration showed in the Figure 1, if both of the two sources begin to send to destination 1 at their peak rate, congestion will occur at the switch. Higher speed links can make the congestion condition in the switch worse.

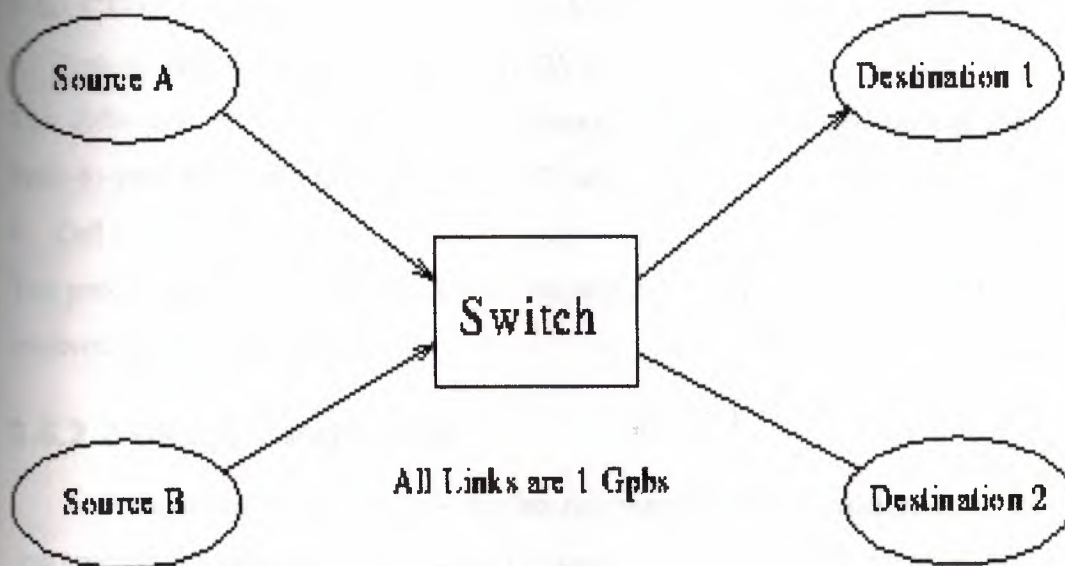


Figure.3.7.A Network with High Speed Links

3. Congestion is caused by slow processors. The problem will be solved when processor speed is improved. This statement can be explained to be wrong similarly to the second one. Faster processors will transmit more data in unit time. If several nodes begin to transmit to one destination simultaneously at their peak rate, the target will be overwhelmed soon. Congestion is a dynamic problem; any static solutions are not sufficient to solve the problem. All the issues presented above: buffer shortage, slow link, slow processor are symptoms not the causes of congestion. Proper congestion management mechanisms is more important than ever.

3.8.2 Connection Parameters

3.8.2.1 Quality Of Service

A set of parameters are negotiated when a connection is set up on ATM networks. These parameters are used to measure the Quality of Service (QOS) of a connection and quantify end-to-end network performance at ATM layer. The network should guarantee the QOS by meet certain values of these parameters.

- Cell Transfer Delay (CTD)}:

The delay experienced by a cell between the first bit of the cell is transmitted by the source and the last bit of the cell is received by the destination. Maximum Cell Transfer Delay (Max CTD) and Mean Cell Transfer Delay (Mean CTD) are used.

- Peak-to-peak Cell Delay Variation (CDV)}:

The difference of the maximum and minimum CTD experienced during the connection. Peak-to-peak CDV and Instantaneous CDV are used.

- Cell Loss Ratio (CLR)}:

The percentage of cells that are lost in the network due to error or congestion and are not received by the destination.

3.8.2.2 Usage Parameters

Another set of parameters is also negotiated when a connection is set up. These parameters discipline the behavior of the user. The network only provide the QOS for the cells that do not violate these specifications.

- Peak Cell Rate (PCR):

The maximum instantaneous rate at which the user will transmit.

- Sustained Cell Rate (SCR):

The average rate as measured over a long interval.

- Burst Tolerance (BT):

The maximum burst size that can be sent at the peak rate.

- Maximum Burst Size (MBS):

The maximum number of back-to-back cells that can be sent at the peak cell rate.

BT and MBS are related as follows: $\text{Burst Tolerance} = (\text{MBS} - 1) \left(\frac{1}{\text{SCR}} - \frac{1}{\text{PCR}} \right)$

- **Minimum Cell Rate (MCR):**

The minimum cell rate desired by a user.

3.8.2.3 Service Categories

Providing desired QOS for different applications is very complex. For example, voice is delay-sensitive but not loss-sensitive, data is loss-sensitive but not delay-sensitive, while some other applications may be both delay-sensitive and loss-sensitive.

To make it easier to manage, the traffic in ATM is divided into five service classes:

- **CBR: Constant Bit Rate**

Quality requirements: constant cell rate, i.e. CTD and CDV are tightly constrained; Low CLR.

Example applications: interactive video and audio.

- **RT-VBR: real-time Variable Bit Rate**

Quality requirements: variable cell rate, with CTD and CDV are tightly constrained; a small nonzero random cell loss is possible as the result of using statistical multiplexing.

Example applications: interactive compressed video.

- **NRT-VBR: Non-Real-Time Variable Bit Rate**

Quality requirements: variable cell rate, with only CTD are tightly constrained; a small nonzero random cell loss is possible as the result of using statistical multiplexing.

Example applications: response time critical transaction processing.

- **UBR: Unspecified Bit Rate**

Quality requirements: using any left over capacity, no CTD or CDV or CLR constrained

Example applications: email and news feed.

- **ABR: Available Bit Rate**

Quality requirements: using the capacity of the network when available and controlling the source rate by feedback to minimize CTD, CDV and CLR.

Example applications: critical data transfer remote procedure call and distributed file services.

These service categories relate traffic characteristics and QOS requirements to network behavior. The QOS requirement for each class is different. The traffic management policy for them are different, too.

The QOS and Usage parameters for these classes are summarized in Table 1:

Table 1: ATM Service Categories

Parameters	CBR	RT-VBR	NRT-VBR	UBR	ABR
PCR and CDVT(5)	specified	specified	specified	specified(3)	specified(4)
SCR, MBS, CDVT(5,6)	N/A	specified	specified	N/A	N/A
MCR(5)	N/A	N/A	N/A	N/A	specified
Peak-to-peak CDV	specified	specified	unspecified	unspecified	unspecified
Mean CTD	unspecified	unspecified	specified	unspecified	unspecified
Max CTD	specified	specified	unspecified	unspecified	unspecified
CLR(5)	specified(1)	specified(1)	specified(1)	unspecified	specified(2)

Notes:

1. The CLR may be unspecified for CLP=1.
2. Minimized for sources that adjust cell flow in response to control information.
3. May not be subject to CAC and UPC procedures.
4. Represents the maximum rate at which the source can send as controlled by the control information.
5. These parameters SRE either explicitly or implicitly specified for PVCs or SVCs.
6. Different values of CDVT may be specified for SCR and PCR.

Among these service classes, ABR is commonly used for data transmissions, which require a guaranteed QOS, such as low probability of loss and error. Small delay is also required for some application, but is not as strict as the requirement of loss and error. Due to the burstiness, unpredictability and huge amount of the data traffic, congestion control of this class is the most needed and is also the most studied.

3.9 What is expected from congestion control

Objectives

The objectives of traffic control and congestion control for ATM are: Support a set of QOS parameters and classes for all ATM services and minimize network and end-system complexity while maximizing network utilization.

3.9.1 Selection Criteria

To design a congestion control scheme is appropriate for ATM network and non-ATM networks as well, the following guidances are of general interest.

- **Scalability**

The scheme should not be limited to a particular range of speed, distance, number of switches, or number of VCs. The scheme should be applicable for both local area networks (LAN) and wide area networks (WAN).

- **Fairness**

In a shared environment, the throughput for a source depends upon the demands by other sources. There are several proposed criterion for what is the correct share of bandwidth for a source in a network environment. And there are ways to evaluate a bandwidth allocation scheme by comparing its results with an optimal result.

- **Fairness Criteria**

1. Max-Min

The available bandwidth is equally shared among connections.

2. MCR plus equal share

The bandwidth allocation for a connection is its MCR plus equal share of the available bandwidth with used MCR removed.

3. Maximum of MCR or Max-Min share

The bandwidth allocation for a connection is its MCR or Max-Min share, which ever is larger.

4. Allocation proportional to MCR

The bandwidth allocation for a connection is weighted proportional to its MCR.

5. Weighted allocation

The bandwidth allocation for a connection is proportional to its pre-determined weight.

- **Fairness Index**

The share of bandwidth for each source should be equal to or converge to the optimal value according to some optimality criterion. We can estimate the fairness of a certain scheme numerically as follows. Suppose a scheme allocates $x_1, x_2, \dots, x_n$, while the optimal allocation is $y_1, y_2, \dots, y_n$. The normalized allocation is $z_i = x_i / y_i$ for each source and the fairness index is defined as following:

$$\text{Fairness} = \frac{\sum(z_i) * \sum(z_i)}{\sum(z_i * z_i)}$$

- **Robustness**

The scheme should be insensitive to minor deviations such as slight mistuning of parameters or loss of control messages. It should also isolate misbehaving users and protect other users from them.

- **Implementability**

The scheme should not dictate a particular switch architecture. It also should not be too complex both in term of time or space it uses.

3.9.2 Generic Functions

It is observed that events responsible for congestion in broadband networks have time constants that differ by orders of magnitude, and multiple controls with appropriate time constants are necessary to manage network congestion.

We can classify the congestion control schemes by the time scale they operate upon: network design, connection admission control (CAC), routing (static or dynamic), traffic shaping, end-to-end feedback control, hop-by-hop feedback control, buffering. The different schemes are functions on different severity of congestion as well as different duration of congestion.

Another classification of congestion control schemes is by the stage that the operation is performed: congestion prevention, congestion avoidance and congestion recovery. Congestion prevention is the method that makes congestion impossible. Congestion avoidance is that the congestion may happen, but the method avoids it by get the network state always in balance. Congestion recovery is the remedy steps to take to pull

the system out of the congestion state as soon as possible and make it less damaging when the congestion already happened.

No matter what kind of scheme is used, the following outstanding problems are the main difficulties that need to be treated carefully: the burstiness of the data traffic, the unpredictability of the resource demand and the large propagation delay versus the large bandwidth.

To meet the objectives of traffic control and congestion control in ATM networks, the following functions and procedures are suggested by the ATM Forum Technical Committee.

3.10 Connection Admission Control

Connection Admission Control (CAC) is defined as the set of actions taken by the network during the call set-up phase in order to determine whether a connection request can be accepted or should be rejected.

Based on the CAC algorithm, a connection request is progressed only when sufficient resources such as bandwidth and buffer space are available along the path of a connection. The decision is made based on the service category, QOS desired and the state of the network which means that the number and conditions of existing connections. Routing and resource allocation are part of CAC when a call is accepted.

3.10.1 Usage Parameter Control

Usage Parameter Control (UPC) is defined as the set of actions taken by the network to monitor and control traffic at the end-system access. Its main purpose is to protect network resources from user misbehavior, which can affect the QOS of other connections, by detecting violations of negotiated parameters and taking appropriate actions.

3.10.2 Generic Cell Rate Algorithm

The Generic Cell Rate Algorithm (GCRA) is used to define conformance with respect to the traffic contract. For each cell arrival, the GCRA determines whether the cells

conforms to traffic contract of the connection. The UPC function may implement GCRA, or one or more equivalent algorithms to enforce conformance.

GCRA is a virtual scheduling algorithm or a continuous-state Leaky Bucket Algorithm as defined by the flowchart in Figure 2 and Figure 3. It is defined with two parameters: the Increment (I) and the Limit (L). The notation GCRA (I,L) is often used.

The GCRA is used to define the relationship between PCR and CDVT, and relationship between SCR and BT. The GCRA is also used to specify the conformance of the declared values of and the above parameters.

3.10.3 Priority Control

The end-system may generate traffic flows of different priority using the Cell Loss Priority (CLP) bit. The network may selectively discard cells with low priority if necessary such as in congestion to protect, as far as possible, the network performance for cells with high priority.

3.10.3.1 Traffic Shaping

Traffic shaping is a mechanism that alters the traffic characteristics of a stream of cells on a connection to achieve better network efficiency whilst meeting the QoS objectives, or to ensure conformance at a subsequent interface.

Examples of traffic shaping are peak cell rate reduction, burst length limiting, reduction of CDV by suitably spacing cells in time, and queue service schemes.

Traffic shaping may be performed in conjunction with suitable UPC functions.

3.10.3.2 Leaky Bucket Algorithm

The most famous algorithm for traffic shaping is leaky bucket algorithm. This method provides a pseudo-buffer (Figure 4). Whenever a user sends a cell, the queue in the pseudo-buffer is increased by one. The pseudo-server serves the queue and the service-time distribution is constant. Thus there are two control parameters in the algorithm: the service rate of the pseudo-server and the pseudo-buffer size.

As long as the queue is not empty, the cells are transmitted with the constant rate of the service rate. So the algorithm can receive a bursty traffic and control the output rate. If



excess traffic makes the pseudo-buffer overflow, the algorithm can choose discarding the cells or tagging them with CLP=1 and transmitting them.

PCR or SCR can be controlled by choosing appropriate values of service rate and buffer size. In addition, combining two buckets with one for each of the parameters can control both PCR and SCR and there are many variances of the original scheme.

3.10.3.3 Network Resource Management

In Network Resource Management (NRM) is responsible for the allocation of network resources in order to separate traffic flows according to different service characteristics, to maintain network performance and to optimize resource utilization. The function is mainly concerned with the management of virtual paths in order to meet QOS requirements.

3.10.3.4 Frame Discard

If a congested network needs to discard cells, it may be better to drop all cells of one frame than to randomly drop cells belonging to different frames, because one cell loss may cause the retransmission of the whole frame, which may cause more traffic when congestion already happened. Thus, frame discard may help avoid congestion collapse and can increase throughput. If done selectively, frame discard may also improve fairness.

3.10.4 Feedback Control

Feedback controls are defined as the set of actions taken by the network and by the end-systems to regulate the traffic submitted on ATM connections according to the state of network elements.

Feedback mechanisms are specified for ABR service class by ATM Forum Technical Committee.

3.10.5 ABR Flow Control

As we have discussed before, the ABR service category uses the link capacity that is left over and is applied to transmit critical data that is sensitive to cell loss. That makes

traffic management for this class the most challenging by the fluctuation of the network load condition, the burstiness of the data traffic itself, and the CLR requirement.

The ATM Forum Technical Committee and Traffic Management Working Group have worked hard on this topic, and here are some of the main issues and the current progress of this area.

3.10.5.1 Some Early Debates

Congestion management in ATM is a hotly debated topic, many contradictory beliefs exist on most issues. These beliefs lead to different approaches in the congestion control schemes. Some of the issues have been closed after a long debate and the ATM Forum Technical Committee finally adopted one of them, and others are still open and the debates are continuing.

3.10.5.1.1 Open-Loop Vs Close-Loop

Open-loop approaches do not need end-to-end feedback, one of the examples of this type are prior-reservation and hop-to-hop flow control. In close-loop approaches, the source adjusts its cell rate in responding to the feedback information received from the network.

It has been argued that close-loop congestion control schemes are too slow in today's high-speed, large range network, by the time a source gets the feedback and reacts to it, several thousand cells may have been lost. But on the other hand, if the congestion has already happened and the overload is of long duration, the condition cannot be released unless the source causing the congestion is asked to reduce its rate. Furthermore, ABR service is designed to use any bandwidth that is left over the source must have some knowledge of what is available when it is sending cells.

The ATM Forum Technical Committee Traffic Management Working Group specified that feedback is necessary for ABR flow control.

3.10.5.1.2 Credit-Based Vs Rate-Based

Credit-Based approaches consist of per-link, per-VC window flow control. The receiver monitors queue lengths of each VC and determines the number of cells the sender

can transmit on that VC, which is called "credit". The sender transmits only as many cells as allowed by the credit.

Rate-Based approaches control the rate by which the source can transmit. If the network is light loaded, the source is allowed to increase its cell rate. If the network is congested, the source should decrease its rate.

After a long debate, ATM Forum finally adopted the rate-based approach and rejected the credit-based approach. The main reason for the credit-based approach not being adopted is that it requires per-VC queuing, which will cause considerable complexity in the large switches, which support millions of VCs. It is not scalable. Rate-Based approaches can work with or without per-VC queuing.

3.10.5.1.3 Binary Feedback Vs Explicit Feedback

Binary Feedback uses one bit in the cell to indicate the elements along the flow path is congested or not. The source will increase or decrease its rate by some predecided rule upon receive the feedback. In Explicit Feedback, the network tells the source exactly what rate is allowed for it to send.

Explicit Rate (ER) feedback approach is preferred, because ER schemes has several advantages over single-bit binary feedback. First, ATM networks are connection oriented and the switches know more information along the flow path, the increased information can only be used by explicit rate feedback. Secondly, the explicit rate feedback is faster to get the source to the optimal operating point. Third, policing is straight forward. The entry switches can monitor the returning message and use the rate directly. Fourth, with fast convergence time, the initial rate has less impact. Fifth, the schemes are robust against errors in or loss of a single message. The next correct message will bring the system to the correct operating point.

There are two ways for explicit rate feedback: forward feedback and backward feedback. With forward feedback, the messages are sent forward along the path and are returned to the source by the destination upon receiving the message. With backward feedback, the messages are sent directly back to the source by the switches whenever congestion condition happens or is pending in any of the switches along the flow path.

3.10.5.1.4 Congestion Detection Queue Length Vs queue Growth

Actually this issue does not cause too much debate. In earlier schemes, large queue length is often used as the indication of congestion. But there are some problems with this method.

First, it is a static measurement. For example, a switch with a 10k cells waiting in queue is not necessarily more congested than a switch with a 10 cell queue if the former one is draining out its queue with 10k cell per second rate and the queue in the latter is building up quickly.

Secondly, using queue length as the method of congestion detection was shown to result in unfairness. Sources that start up late were found to get lower throughput than those which start early.

Queue growth rate is more appropriate as the parameter to monitor the congestion state because it shows the direction that the network state is going. It is natural and direct to use queue growth rate in a rate-based scheme, with the controlled parameter and the input parameter have the same unit.

3.11 RM-cell Structure

In the ABR service, the source adapts its rate to changing network conditions. Information about the state of the network like bandwidth availability, state of congestion, and impending congestion, is conveyed to the source through special control cells called Resource Management Cells (RM-cells).

ATM Forum Technical Committee specifies the format of the RM-cell. The already defined fields in a RM-cell that is used in ABR service is explained in this section.

1. Header

The first five bytes of a RM-cell are the standard ATM header with PTI=110 for a VCC and VCI=6 for a VPC.

2. ID

The protocol ID. The ITU has assigned this field to be set to 1 for ABR service.

3. **DIR**

Direction of the RM-cell with respect to the data flow which it is associated with. It is set to 0 for forward RM-cells and 1 for backward RM-cells.

4. **BN**

Backward Notification. It is set to 1 for switch generated (BECN) RM-cells and 0 for source generated RM-cells.

5. **CI**

Congestion Indication. It is set to 1 to indicate congestion and 0 otherwise.

6. **NI**

No Increase. It is set to 1 to indicate no additive increase of rate allowed when a switch senses impending congestion and 0 otherwise.

7. **ER**

Explicit rate. It is used to limit the source rate to a specific value.

8. **CCR**

Current Cell Rate. It is used to indicate to current cell rate of the source.

9. **MCR**

Minimum Cell Rate. The minimum cell rate desired by the source.

3.11.1 Service Parameters

ATM Forum Technical Committee defined a set of flow control parameters for ABR service.

1. **PCR**

Peak Cell Rate, it is the source desired but the maximum rate the network can support. It is negotiated when the connection is set up.

2. **MCR**

Minimum Cell Rate, the source need not to reduce its rate below it under any condition. It is negotiated when the connection is set up.

3. **ICR**

Initial Cell Rate, the startup rate after idle periods. It is negotiated when the connection is set up.

4. AIR

Additive Increase Rate, the highest rate increase possible. It is negotiated when the connection is set up.

5. Nrm

The number of cells transmitted per RM-cell sent. It is negotiated when the connection is set up.

6. Mrm

Used by the destination to control allocation of bandwidth between forward RM-cells, backward RM-cells, and data cells. It is negotiated when the connection is set up.

7. RDF

Rate Decrease Factor, to control the number of cells sent upon idle startup before the network can establish control in one Round Trip Time (RTT). It is negotiated when the connection is set up.

8. ACR

Allowed Cell Rate, the source can not transmit with rate higher than it.

9. Xrm

The maximum RM-cells sent without feedback before the source need to reduce its rate. It is negotiated when the connection is set up.

10. TOF

Time Out Factor, to control the maximum time permitted between sending forward RM-cells before a rate decrease is required. It is negotiated when the connection is set up.

11. Trm

The inter-RM time interval used in the source behavior. It is negotiated when the connection is set up.

12. RTT

Round Trip Time between the source and the destination. It is computed during call setup.

13. XDF

Xrm Decrease Factor, specify how much of the reduction of the source rate when XRM is triggered. It is negotiated when the connection is set up.

These parameters are used to implement ABR flow-control on a per-connection basis, and the source, switch and destination must behave within the rules that defined by these parameters.

The function and usage of these parameters are still under study.

3.11.2 Source, Destination and Switch Behavior

ATM Forum Technical Committee also specifies the source, destination, and switch behavior for the service.

There are two notations that need to be explained before we discuss the network behavior.

In-Rate Cells: The cells that counted in the user's rate with $CLP=0$. In-rate cells include data cells and in-rate RM-cells.

Out-of-Rate Cells: These cells are RM-cells and are not counted in the user's rate. They are used when $ACR=0$ and in-rate RM cells can not be send. The CLP is set to 1 for them.

In this section, we discuss some highlights of the specification.

- **Source Behavior**

1. The value of ACR shall never exceed PCR , nor shall it ever be less than MCR .

The source shall never send in-rate cells at a rate exceeding ACR .

2. The source shall start with ACR at ICR and the first in-rate cell sent shall be a forward RM-cell.

3. The source shall send one RM-cell after every N_{rm} data cells.

4. If the source does not receive any feedback since it sends the last RM-cell, it shall reduce its rate by at least $ACR \cdot T \cdot TDF$ after $TOF \cdot N_{rm}$ cell intervals.

5. If at least X_{rm} in-rate forward RM-cells have been sent since the last backward RM-cell with $BN=0$ was received, ACR shall be reduced by at least $ACR \cdot XDF$.

6. When a backward RM-cell is received with $CI=1$, ACR shall be reduced by at least $ACR \cdot N_{rm} / RDF$.

If the backward RM-cell has both $CI=0$ and $NI=0$, the ACR may be increased by no more than $AIR \cdot N_{rm}$.

7. Out-of-rate forward RM-cells shall not be sent at a rate greater than TCR .

- **Destination Behavior**

1. When a data cell is received, the destination shall save the EFCI state.

2. When returning a RM-cell, it shall set CI if saved EFCI is 1.

Congested destination may set both CI and NI, or reduce ER.

3. If a RM-cell has not been returned while the next one arrives, throw away the old one.

4. The destination can generate a backward RM-cell without having received a forward RM-cell.

- **Switch Behavior**

1. The switch may set the EFCI flag in the data cell headers.

2. The switch may set CI or NI in the RM-cells, or may reduce the ER field.

3. The switch may generate backward RM-cells with CI or NI set.

3.12 Representative Schemes

The following is a brief description of congestion control schemes that are proposed to the ATM Forum. The various mechanisms can be classified broadly depending upon the congestion monitoring criteria used and the feedback mechanism employed.

- **EFCI Control Schemes**

These classes of feedback mechanisms use binary feedback involving the setting of the EFCI bit in the cell header. The simplest example of a binary feedback mechanism is based on the old DEC bit scheme. In this scheme all the VCs in a switch share a common FIFO queue and the queue length is monitored. When the queue length exceeds a threshold, congestion is declared and the cells passing the switch have their EFCI bit set. When the queue length falls below the threshold the cells are passed without their EFCI bit set. The source will adjust its rate accordingly when it sees the feedback cells with the EFCI bit set or not. Variations of this scheme include using two thresholds for the indication and removing congestion respectively.

Binary feedback mechanisms can sometimes be fair because long hop VCs have higher possibility to have their cell EFCI bit set and get fewer opportunities to increase their rate. It is called the "beat down problem". This problem can be alleviated by some enhancements to the basic scheme such provide separate queues for each VCs.

But a coherent problem with binary feedback mechanisms are that they are too slow for rate-based control in high-speed networks.

- **Explicit Rate Feedback Schemes**

As we have discussed before, explicit rate feedback control would not only be faster but would offer more flexibility to switch designers. Many explicit rate feedback control schemes have been proposed, the following are some that are documented by the ATM Forum.

1. **Enhanced Proportional Rate Control Algorithm (EPRCA)**

In EPRCA, the source sends data cells with EFCI set to 0 and sends RM-cells every n data cells. The RM-cells contain desired explicit rate (ER), current cell rate (CCR) and congestion indication (CI). The source usually initializes CCR to the allowed cell rate (ACR) and CI to zero.

The switch computes a mean allowed cell rate (MACR) for all VCs using exponential weighted average:

$$\text{MACR} = (1 - \alpha) * \text{MACR} + \alpha * \text{CCR}$$

and the fair share as a fraction of this average, where α and the fraction are chosen to be $1/16$ and $7/8$ respectively. The ER field in the returning RM-cells are reduced to fair share if necessary. The switch may also set the CI bit in the cells passing when it is congested which is sensed by monitoring its queue length.

The destination monitors the EFCI bits in data cells and marks the CI bit in the RM-cell if the last seen data cell had EFCI bit set.

The source decreases its rate continuously after every cell by a fixed factor and increases its rate by a fixed amount if the CI bit is not set. Another rule is that the new increased rate must never exceed either the ER in the returned cell or the PCR of the connection. The main problem of this scheme is that the congestion detection is based on the queue-length and this method is shown to result in unfairness. Sources that start up late may get lower throughput than those that start early.

2. **Target Utilization Band (TUB) Congestion Avoidance Scheme**

In each switch, a target rate is defined as slightly below the link bandwidth, such as 85-90% of the full capacity. The input rate of the switch is measured over a fixed averaging interval. The load factor z is then computed as:

$$\text{Load Factor } z = \text{Input Rate} / \text{Target Rate}$$

When the load factor is far from z , which means the switch is either highly overloaded or highly underloaded, all VCs are asked to change their load by this factor z . When the load factor is close to 1, between $1-\delta$ and $1+\delta$ for a small δ , the switch gives different feedback to underloading sources and overloading sources.

A fairshare is computed, and all sources whose rates are more than the fair share are asked to divide their rates by $z/(1+\delta)$, while those below the fair share are asked to divide their rates by $z/(1-\delta)$.

3. Explicit Rate Indication for Congestion Avoidance (ERICA)

This scheme tries to achieve efficiency and fairness concurrently by allowing underloaded VCs to increase their rate to fair share in spite of the conditions of the network and the sources already equal or greater than fair share may increase their rate if the link is under used. And the target capacity of the switch is set higher, 90-95% of the full bandwidth.

The switch calculates fair share as:

Fairshare = Target capacity / Number of active VCs

And the remaining capacity that a source can use is:

VCshare = CCR / Load Factor z

Then the switch sets the source's rate to the maximum of the two.

The information used to compute the quantities comes from the forward RM-cells and the feedback is given in the backward RM-cells. This ensures that the most current information is used to provide fastest feedback.

Another advantage of this scheme is that it has few parameters which can be tuned easily.

4. Congestion Avoidance Using Proportional Control (CAPC)

In this scheme, the switches also set a target utilization slightly below 1 and measure the input rate to compute load factor z .

During underload (z is less than 1), fair share is increased as:

Fairshare = Fairshare * Min(ERU, $1+(1-z)*Rup$)

where Rup is a slope parameter between 0.025 and 0.1, and ERU is the maximum increase allowed.

During overload (z is greater than 1), fair share is decreased as:

Fairshare = Fairshare * Max(ERF, $1-(z-1)*Rdn$)

where R_{dn} is a slope parameter between 0.2 and 0.8, and ERF is the minimum decrease required.

The source should never allowed to transmit at a rate higher than the fair share.

The distinguishing feature of this scheme is that it is oscillation-free in steady state.

5. ER Based on Bandwidth Demand Estimate Algorithm

The switch calculates the Mean Allowed Cell rate (MACR) basing on a running exponential average of the ACR value from each VC's forward RM-cells as:

$$MACR = MACR + (ACR - MACR) * AVF$$

where AVF (ACR Variation Factor) is set to $1/16$.

If the load factor is less than 1, the left-over bandwidth is reallocated according to:

$$MACR = MACR + MAIR$$

where $MAIR$ is the MACR Additive Increase Rate.

The ER value is computed as:

$$ER = MACR * MRF$$

where MRF is the MACR Reduction Factor if congestion is detected,

$$ER = MACR$$

if no congestion is detected.

The congestion condition is detected by observing that the queue derivative is positive.

CHAPTER 4

A Survey of ATM Switching Techniques

Overview

Asynchronous Transfer Mode (ATM) switching is not defined in the ATM standards, but a lot of research has been done to explore various ATM switch design alternatives. Each design has its own merits and drawbacks, in terms of throughput, delay, scalability, buffer sharing and fault tolerance. By examining the features of the basic switch designs, several conclusions can be inferred about the design principles of ATM switching, and the various tradeoffs involved in selecting the appropriate approach.

Switching

Asynchronous Transfer Mode (ATM) is the technology of choice for the Broadband Integrated Services Digital Network (B-ISDN). The ATM is proposed to transport a wide variety of services in a seamless manner. In this mode, user information is transferred in a connection-oriented fashion between communicating entities using fixed-size packets, known as ATM cells. The ATM cell is fifty-three bytes long, consisting of a five byte header and a forty-eight byte information field, sometimes referred to as payload.

Because switching is not part of the ATM standards, vendors use a wide variety of techniques to build their switches. A lot of research has been done to explore the different switch design alternatives, and if one were to attempt to describe the different variations of each, this paper would turn into a book of several volumes. Hence only the major issues are high-lighted here, and the various merits and drawbacks of each approach are briefly examined.

The aim of ATM switch design is to increase speed, capacity and overall performance. ATM switching differs from conventional switching because of the high-speed interfaces (50 Mbps to 2.4 Gbps) to the switch, with switching rates up to 80 Gbps in the backplane. In addition, the statistical capability of the ATM streams passing through the ATM switching systems places additional demands on the switch. Finally, transporting various types of traffic, each with different requirements of behavior and semantic or time

transparency (cell loss, errors, delays) is not a trivial matter. To meet all these requirements, the ATM switches had to be significantly different from conventional switches.

A large number of ATM switch design alternatives exist, both in the hardware and the software components. The software needs to be partitioned into hardware-dependent and independent components, and the modularity of the whole switch design is essential to easily upgrade the switches. An ATM switching system is much more than a fabric that simply routes and buffers cells (as is usually meant by an ATM switch), rather it comprises an integrated set of modules. Switching systems not only relay cells, but also perform control and management functions. Moreover, they must support a set of traffic control requirements.

Due to this functional division of an ATM switching system, the remainder of this survey is organized as follows. First, various switching functions and requirements are discussed. Then a generic functional model for a switching architecture is presented to simplify the ensuing discussion. The core of the switch, the switch fabric, will be the main focus of this paper, and thus it is explored in great depth, highlighting the main design categories. Finally, the problems and tradeoffs exposed when analyzing these switch fabric design alternatives are used to draw some conclusions on the major switch design principles.

4.1 Switching Functions

An ATM switch contains a set of input ports and output ports, through which it is interconnected to users, other switches, and other network elements. It might also have other interfaces to exchange control and management information with special purpose networks. Theoretically, the switch is only assumed to perform cell relay and support of control and management functions. However, in practice, it performs some internetworking functions to support services such as SMDS or frame relay. It is useful to examine the switching functions in the context of the three planes of the B-ISDN model.

4.1.1 User Plane

The main function of an ATM switch is to relay user data cells from input ports to the appropriate output ports. The switch processes only the cell headers and the payload is carried transparently. As soon as the cell comes in through the input port, the Virtual Path Identifier / Virtual Channel Identifier (VPI/VCI) information is derived and used to route the cells to the appropriate output ports. This function can be divided into three functional blocks: the input module at the input port, the cell switch fabric (sometimes referred to as switch matrix) that performs the actual routing, and the output modules at the output ports.

4.1.2 Control Plane

This plane represents functions related to the establishment and control of the VP/VC connections. Unlike the user data cells, information in the control cells payload is not transparent to the network. The switch identifies signaling cells, and even generates some itself. The Connection Admission Control (CAC) carries out the major signaling functions required. Signaling information may/may not pass through the cell switch fabric, or maybe exchanged through a signaling network such as SS7.

4.1.3 Management Plane

The management plane is concerned with monitoring the controlling the network to ensure its correct and efficient operation. These operations can be subdivided as fault management functions, performance management functions, configuration management functions, security management functions, accounting management functions and traffic management functions. These functions can be represented as being performed by the functional block Switch Management.

The Switch Management is responsible for supporting the ATM layer Operations and Maintenance (OAM) procedures. OAM cells may be recognized and processed by the ATM switch. The switch must identify and process OAM cells, maybe resulting in generating OAM cells. As with signaling cells, OAM cells may/may not pass through cell switch fabric. Switch Management also supports the interim local management interface (ILMI) of

the UNI. The Switch Management contains, for each UNI, a UNI management entity (UME), which may use SNMP.

4.1.4 Traffic Control Functions

The switching system may support connection admission control, usage/network parameter control (UPC/NPC), and congestion control. We will regard UPC/NPC functions as handled by the input modules, congestion control functions as handled by the Switch Management, while special buffer management actions (such as cell scheduling and discarding) are supervised by the Switch Management, but performed inside the cell switch fabric where the buffers are located. Section 5.7 will examine how buffer management is carried out.

4.2 A Generic Switching Architecture

It will be useful to adopt a functional block model to simplify the discussion of various design alternatives. Throughout this chapter, we will divide switch functions among the previously defined broad functional blocks: input modules, output modules, cell switch fabric, connection admission control, and Switch Management. Figure 4.1 illustrates this switching model. These functional blocks are service-independent, and the partitioning does not always have well-defined boundaries between the functional blocks.

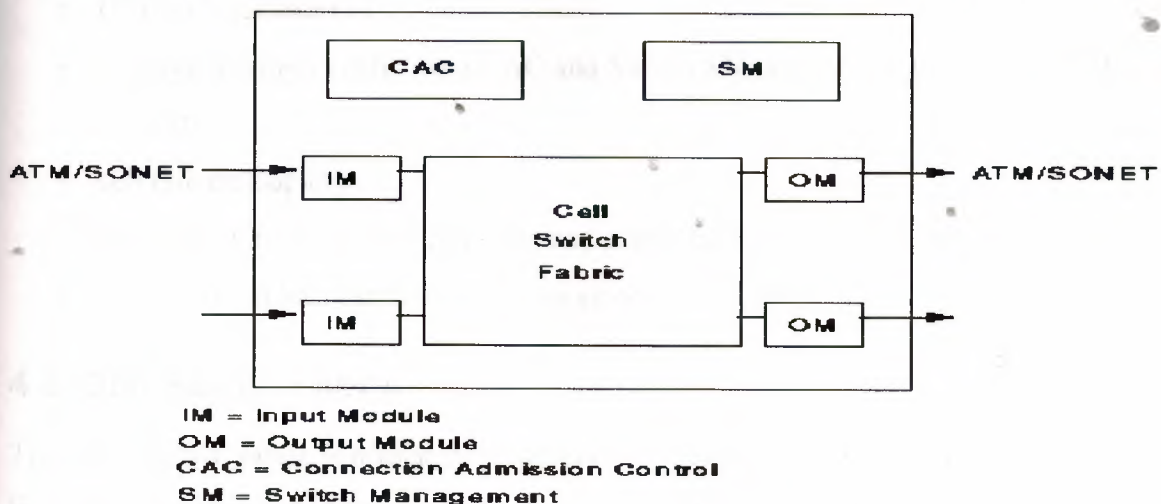


Figure 4.1: A generic switching model (adapted from Chen and Liu [3])

4.3 Switch Interface

4.3.1 Input Modules

The input module first terminates the incoming signal (assume it is a SONET signal) and extracts the ATM cell stream. This involves signal conversion and recovery, processing SONET overhead, and cell delineation and rate decoupling. After that, for each ATM cell the following functions should be performed:

- error checking the header using the Header Error Control (HEC) field
- validation and translation of VPI/VCI values
- determination of the destination output port
- passing signaling cells to CAC and OAM cells to Switch Management
- UPC/UNC for each VPC/VCC
- addition of an internal tag containing internal routing and performance monitoring information for use only within the switch

4.3.2 Output Modules

These prepare the ATM cell streams for physical transmission by:

- removing and processing the internal tag
- possible translation of VPI/VCI values
- HEC field generation
- possible mixing of cells from CAC and Switch Management with outgoing cell streams
- cell rate decoupling
- mapping cells to SONET payloads and generation of SONET overhead
- conversion of the digital bitstream to an optical signal

4.4 Cell Switch Fabric

The cell switch fabric is primarily responsible for routing of data cells and possibly signaling and management cells as well. Since the remainder of this paper focuses on the cell

switch fabric, the next section is devoted to exploring its various components in considerable detail.

4.5 Connection Admission Control (CAC)

Establishes, modifies and terminates virtual path/channel connections. More specifically, it is responsible for:

- high-layer signaling protocols
- signaling ATM Adaptation Layer (AAL) functions to interpret or generate signaling cells
- interface with a signaling network
- negotiation of traffic contracts with users requesting new VPCs/VCCs
- renegotiation with users to change established VPCs/VCCs
- allocation of switch resources for VPCs/VCCs, including route selection
- admission/rejection decisions for requested VPCs/VCCs
- generation of UPC/NPC parameters

If the CAC is centralized, a single processing unit would receive signaling cells from the input modules, interpret them, and perform admission decisions and resource allocation decisions for all the connections in the switch. CAC functions may be distributed to blocks of input modules where each CAC has a smaller number of input ports. This is much harder to implement, but solves the connection control processing bottleneck problem for large switch sizes, by dividing this job to be performed by parallel CACs. A lot of information must be communicated and coordinated among the various CACs. In Hitachi's and NEC's ATM switches, input modules- each with its CAC- also contain a small ATM routing fabric. Some of the distributed CAC functions can also be distributed among output modules which can handle encapsulation of high-layer control information into outgoing signaling cells.

4.6 Switch Management

Handles physical layer OAM, ATM layer OAM, configuration management of switch components, security control for the switch database, usage measurements of the switch

resources, traffic management, administration of a management information base, customer-network management, interface with operations systems and finally support of network management. This area is still under development, so the standards are not established yet. Switch Management is difficult because management covers an extremely wide spectrum of activities. In addition, the level of management functions implemented in the switch can vary between minimal and complex. Switch Management must perform a few basic tasks. It must carry out specific management responsibilities, collect and administer management information, communicate with users and network managers, and supervise and coordinate all management activities. Management functions include fault management, performance management, configuration management, accounting management, security management, and traffic management. Carrying out these functions entails a lot of intraswitch communication between the Switch Management and other functional blocks.

A centralized Switch Management can be a performance bottleneck if it is overloaded by processing demands. Hence, Switch Management functions can be distributed among input modules, but a lot of coordination would be required. Each distributed input module Switch Management unit can monitor the incoming user data cell streams to perform accounting and performance measurement. Output module Switch Management units can also monitor outgoing cell streams.

4.7 The Cell Switch Fabric

The cell switch fabric is primarily responsible for transferring cells between the other functional blocks (routing of data cells and possibly signaling and management cells as well). Other possible functions include:

- cell buffering
- traffic concentration and multiplexing
- redundancy for fault tolerance
- multicasting or broadcasting
- cell scheduling based on delay priorities
- congestion monitoring and activation of Explicit Forward Congestion Indication (EFCI)

Each of these functions will be explored in depth in the context of the various design alternatives and principles.

4.7.1 Concentration, Expansion and Multiplexing

Traffic needs to be concentrated at the inputs of the switching fabric to better utilize the incoming link connected to the switch. The concentrator aggregates the lower variable bit rate traffic into higher bit rate for the switching matrix to perform the switch at standard interface speed. The concentration ratio is highly correlated with the traffic characteristics, so it needs to be dynamically configured. The concentrator can also aid in dynamic traffic distribution to multiple routing and buffering planes, and duplication of traffic for fault tolerance. At the outputs of the routing and buffering fabric, traffic can be expanded and redundant traffic can be combined.

4.7.2 Routing and Buffering

The routing and buffering functions are the two major functions performed by the cell switch fabric. The input module attaches a routing tag to each cell, and the switch fabric simply routes the arriving cells from its inputs to the appropriate outputs. Arriving cells may be aligned in time by means of single-cell buffers. Because cells may be addressed to the same output simultaneously, buffers are needed. Several routing and buffering switch designs have aided in setting the important switch design principles. All current approaches employ a high degree of parallelism, distributed control, and the routing function is performed at the hardware level.

Before examining the impact of the various design alternatives, we need to consider the essential criteria for comparing among them. The basic factors are:

1. throughput (total output traffic rate/input traffic rate)
2. utilization (average input traffic rate/maximum possible output traffic rate)
3. cell loss rate
4. cell delays
5. amount of buffering
6. complexity of implementation

Traditionally switching has been defined to encompass either space switching or time switching or combinations of both techniques. The classification adopted here is slightly different in the sense that it divides the design approaches under the following four broad categories:

1. shared memory
2. shared medium
3. fully interconnected
4. space division

For simplicity, the ensuing discussion will assume a switch with N input ports, N output ports, and all port speeds equal to V cells/s. Multicasting and broadcasting will be addressed with the other issues in the next section, so they will be temporarily ignored in this discussion.

4.7.2.1 Shared Memory Approach

Figure 4.2 illustrates the basic structure of a shared memory switch. Here incoming cells are converted from serial to parallel form, and written sequentially to a dual port Random Access Memory. A memory controller decides the order in which cells are read out of the memory, based on the cell headers with internal routing tags. Outgoing cells are demultiplexed to the outputs and converted from parallel to serial form.

This approach is an output queuing approach, where the output buffers all physically belong to a common buffer pool. The approach is attractive because it achieves 100% throughput under heavy load. The buffer sharing minimizes the amount of buffers needed to achieve a specified cell loss rate. This is because if a large burst of traffic is directed to one output port, the shared memory can absorb as much as possible of it. CNET's Prelude switch was one of the earliest prototypes of this technique, which employed slotted operation with packet queuing. Hitachi's shared buffer switch and AT&T's GCNS-2000 are famous examples of this scheme.

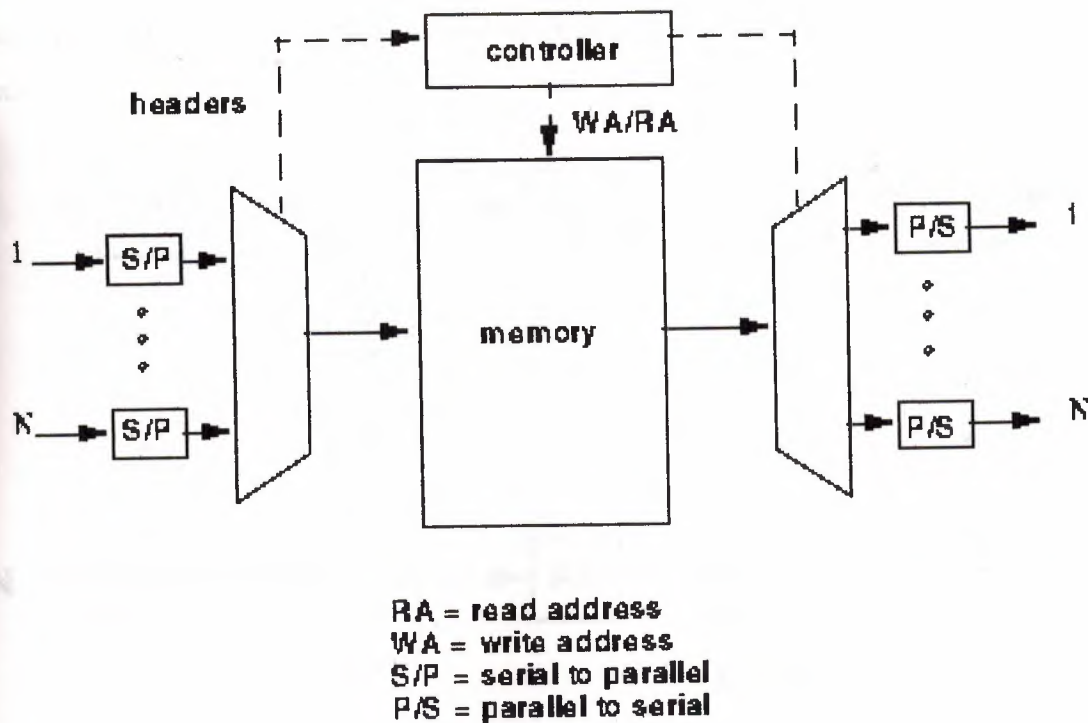


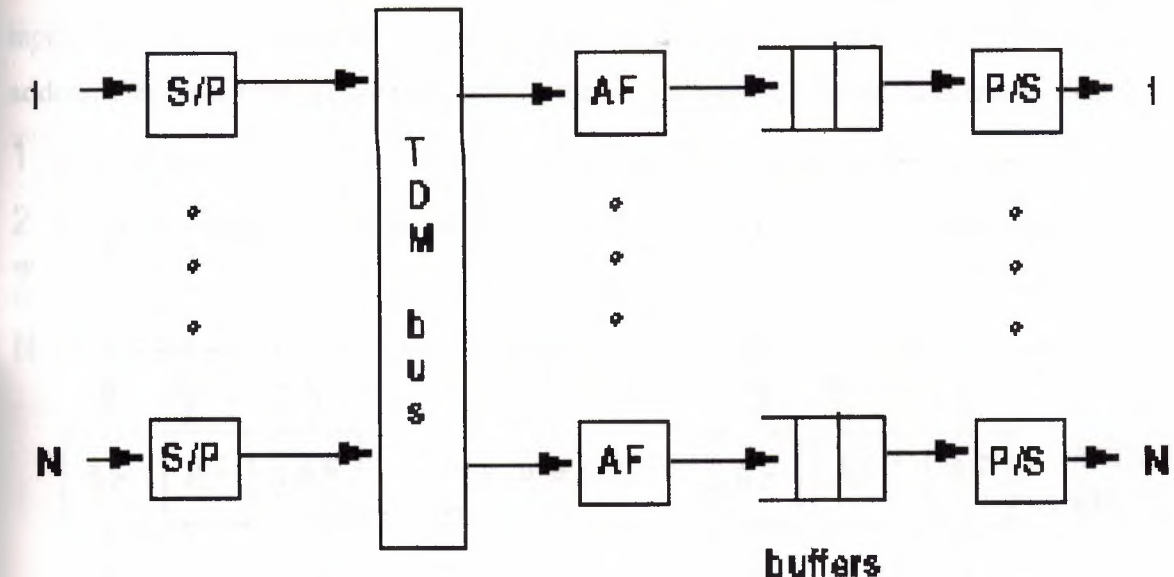
Figure 4.2: Basic structure of a shared-memory switch

The approach, however, suffers from a few drawbacks. The shared memory must operate N times faster than the port speed because cells must be read and written one at a time. As the access time of memory is physically limited, the approach is not very scalable. The product of the number of ports times port speed (NV) is limited. In addition, the centralized memory controller must process cell headers and routing tags at the same rate as the memory. This is difficult for multiple priority classes, complicated cell scheduling, multicasting and broadcasting.

4.7.2.2 Shared Medium Approach

Cells may be routed through a shared medium, like a ring, bus or dual bus. Time-division multiplexed buses are a popular example of this approach, and figure 3 illustrates their structure. Arriving cells are sequentially broadcast on the TDM bus in a round-robin manner. At each output, address filters pass the appropriate cells to the output buffers,

based on their routing tag. The bus speed must be at least NV for cells/s to eliminate input queuing.



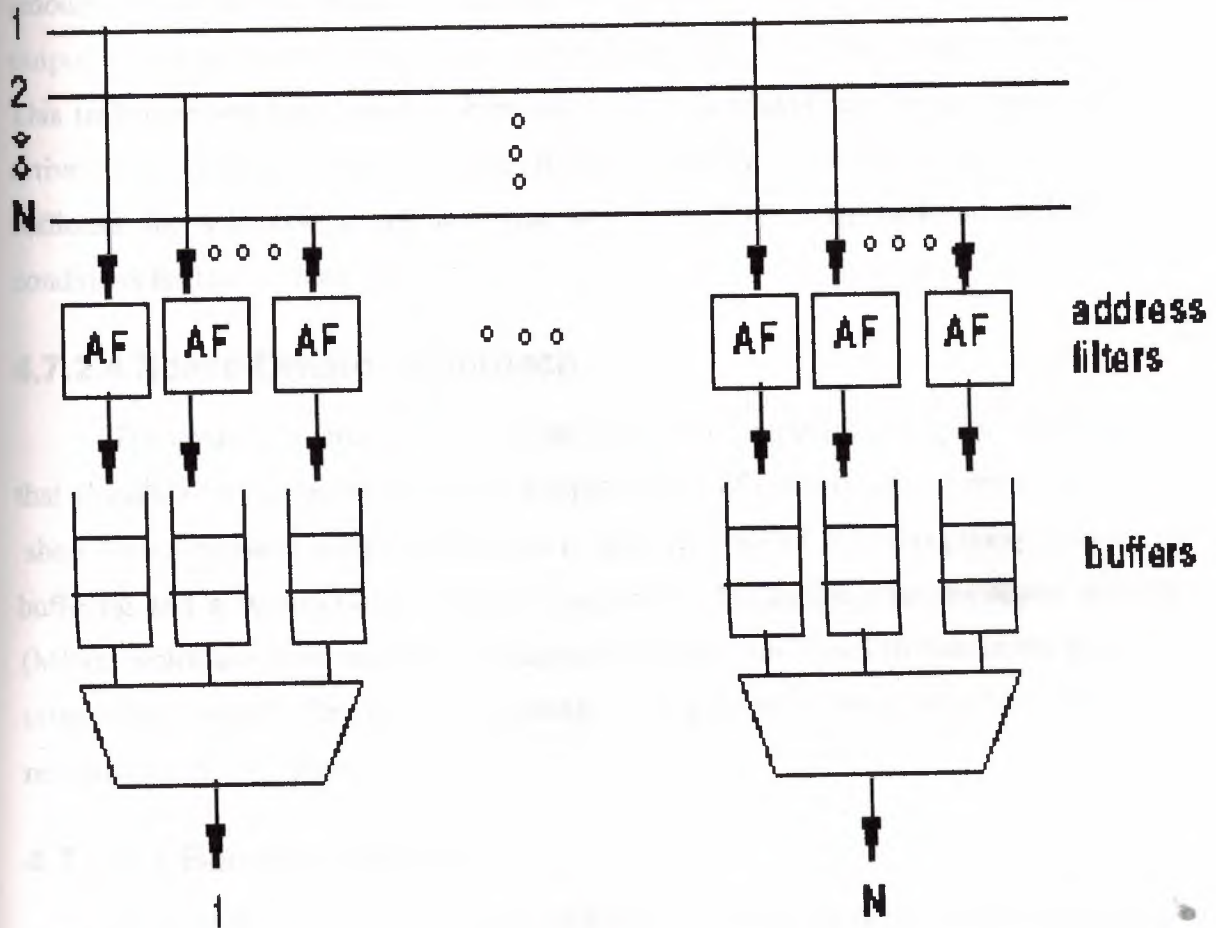
AF = address filter
S/P = serial to parallel
P/S = parallel to serial

Figure 4.3: A shared bus switch

The outputs are modular, which makes address filters and output buffers easy to implement. Also the broadcast-and-select nature of the approach makes multicasting and broadcasting straightforward. As a result, many such switches have been implemented, such as IBM's Packetized Automated Routing Integrated System (PARIS) and plaNET, NEC's ATM Out-put Buffer Modular Switch (ATOM), and Fore Systems' ForeRunner ASX-100 to mention a few. The Synchronous Composite Packet Switching (SCPS) which uses multiple rings is also one of the most famous experiments of shared medium switches. However, because the address filters and output buffers must operate at the shared medium speed, which is N times faster than the port speed, this places a physical limitation on the scalability of the approach. In addition, unlike the shared memory approach, output buffers are not shared, which requires more total amount of buffers for the same cell loss rate.

4.7.2.3 Fully Interconnected Approach

In this approach, independent paths exist between all N squared possible pairs of inputs and outputs. Hence arriving cells are broadcast on separate buses to all outputs and address filters pass the appropriate cells to the output queues. This architecture is illustrated



in fig 4.

Figure 4.4: A fully interconnected switch

This design has many advantages. As before, all queuing occurs at the outputs. In addition, multicasting and broadcasting are natural, like in the shared medium approach. Address filters and output buffers are simple to implement and only need to operate at the port speed. Since all of the hardware operates at the same speed, the approach is scalable to any size and speed. Fujitsu's bus matrix switch and GTE Government System's SPANet are examples of switches in which this design was adopted. Unfortunately, the quadratic

growth of buffers limits the number of output ports for practical reasons. However, the port speed is not limited except by the physical limitation on the speed of the address filters and output buffers.

The Knockout switch developed by AT&T was an early prototype where the amount of buffers was reduced at the cost of higher cell loss. Instead of N buffers at each output, it was proposed to use only a fixed number of buffers L for a total of $N \times L$ buffers. This technique was based on the observation that it is unlikely that more than L cells will arrive for any output at the same time. It was argued that selecting the L value of 8 was sufficient for achieving a cell loss rate of $1/1$ Million under uniform random traffic conditions for large values of N .

4.7.2.4 Space Division Approach

The crossbar switch is the simplest example of a matrix-like space division fabric that physically interconnects any of the N inputs to any of the N outputs. Genda et al.

show how a crossbar switch can be used to achieve a rate of 160 Gbps, using input/output buffering and a bi-directional arbitration algorithm. Multistage interconnection networks (MINs) which are more tree-like structures, were then developed to reduce the N squared crosspoints needed for circuit switching, multiprocessor interconnection and, more recently, packet switching.

4.7.2.4.1 Banyan networks

One of the most common types of MINs is the Banyan network (it is interesting to note that it was given this name because its shape resembles the tropical tree of that name). The Banyan network is constructed of an interconnection of stages of switching elements. A basic 2×2 switching element can route an incoming cell according to a control bit (output address). If the control bit is 0, the cell is routed to the upper port address; otherwise it is routed to the lower port address. To better understand the composition of Banyan networks, consider forming a 4×4 Banyan network. Figure 5 shows the step-by-step interconnection of switching elements to form 4×4 , and then 8×8 Banyan networks. The interconnection of two stages of 2×2 switching elements can be done by using the first bit of the output

address to denote which switching element to route to, and then using the last bit to specify the port.

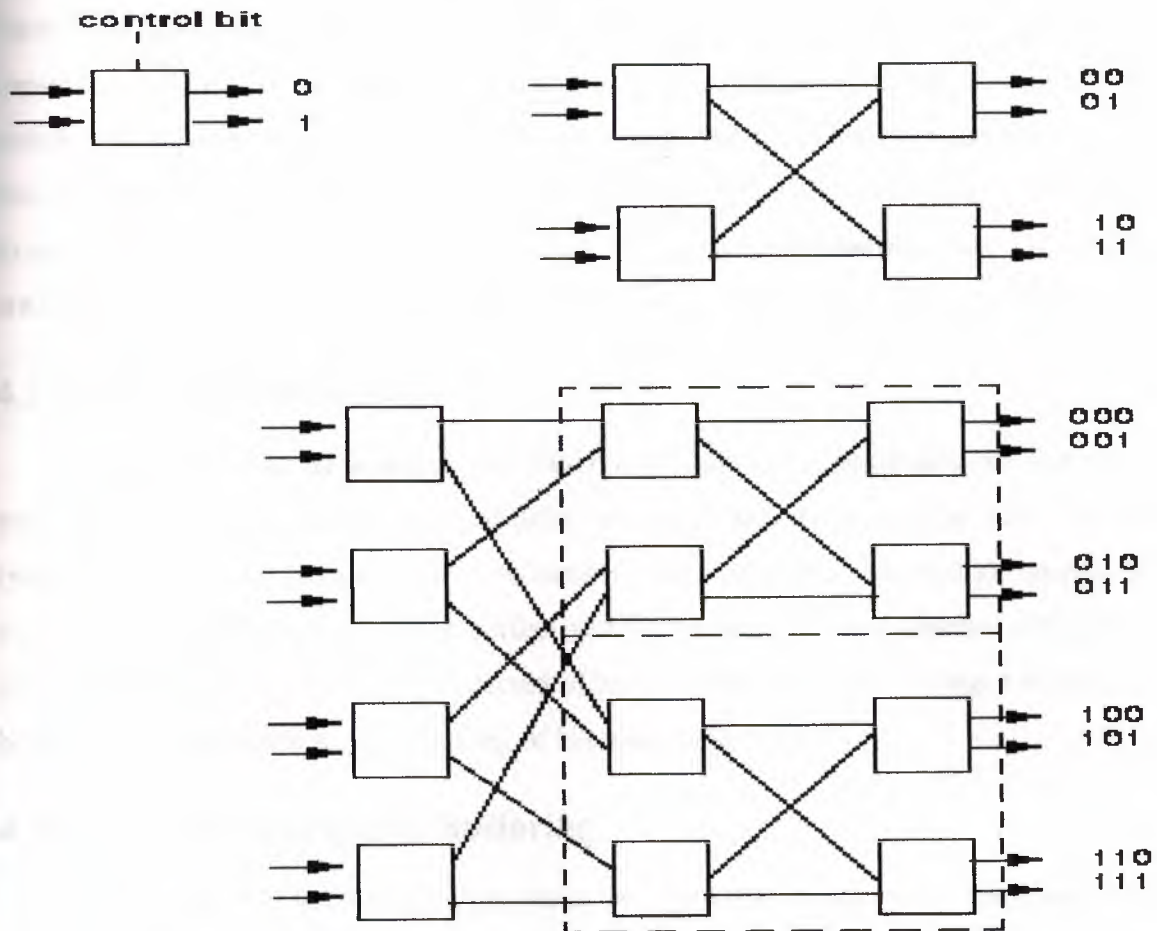


Figure 4. 5: Switching element, 4x4 Banyan network and 8x8 Banyan network

8x8 bananas can be recursively formed by using the first bit to route the cell through the first stage, either to the upper or lower 4x4 network, and then using the last 2 bits to route the cell through the 4x4 network to the appropriate output port. In general, to construct a $N \times N$ Banyan network, the n th stage uses the n th bit of the output address to route the cell. For $N = 2$ to the power of n , the Banyan will consist of $n = \log$ to the base 2 of N stages, each consisting of $N/2$ switching elements. A MIN is called self-routing when the output address completely specifies the route through the network (also called digit-controlled routing).

The Banyan network technique is popular because switching is performed by simple switching elements, cells are routed in parallel, all elements operate at the same speed (so there is no additional restriction on the size N or speed V), and large switches can be easily constructed modularly and recursively and implemented in hardware. Bellcore's Sunshine switch, and Alcatel Data Networks' 1100 are just a few examples of switches employing this technique. It is clear that in a Banyan network, there is exactly one path from any input to any output. Regular Banyans use only one type of switching element, and SW-Banyans are a subset of regular Banyans, constructed recursively from $L \times M$ switching elements.

4.7.2.4.2 Delta Networks

Delta networks are a subclass of SW-Banyan networks, possessing the self-routing property. There are numerous types of delta networks, such as rectangular delta networks (where the switching elements have the same number of outputs as inputs), omega, flip, cube, shuffle-exchange (based on a perfect shuffle permutation) and baseline networks. A delta- b network of size $N \times N$ is constructed of $b \times b$ switching elements arranged in $\log$ to the base b of N stages, each stage consisting of N/b switching elements.

4.7.2.4.2.1 Blocking and Buffering

Unfortunately, since Banyan networks have less than N squared crosspoints, routes of two cells addressed to two different outputs might conflict before the last stage. When this situation, called internal blocking, occurs, only one of the two cells contending for a link can be passed to the next stage, so overall throughput is reduced. A solution to this problem is to add a sort network (such as a Batcher bitonic sort network) to arrange the cells before the Banyan network. This will be internally non-blocking for cells addressed to different outputs. However, if cells are addressed to the same output at the same time, the only solution to the problem is buffering. Buffers can be placed at the input of the Batcher network, but this can cause "head-of-line" blocking, where cells wait for a delayed cell at the head of the queue to go through, even if their own destination output ports are free. This situation can be remedied by First-In-First-Out buffers, but these are quite complex to implement.

Alternatively, buffers may be placed internally within the Banyan switching elements. Thus if two cells simultaneously attempt to go to the same output link, one of them is buffered within the switching element. This internal buffering can also be used to implement a backpressure control mechanism, where queues in one stage of the Banyan will hold up cells in the preceding stage by a feedback signal. The backpressure may eventually reach the first stage, and create queues at the Banyan network inputs. It is important to observe that internal buffering can cause head-of-line blocking at each switching element, and hence it does not achieve full throughput. Awdeh and Mouftah have designed a delta-based ATM switch with backpressure mechanism capable of achieving a high throughput, while significantly reducing the overall required memory size.

A third alternative is to use a recirculating buffer external to the switch fabric. This technique has been adopted in Bellcore's Sunshine and AT&T's Starlite wideband digital switch. Here output conflicts are detected after the Batchier sorter, and a trap network selects a cell to go through, and recirculates the others back to the inputs of the Batchier network. Unfortunately, this approach requires complicated priority control to maintain the sequential order of cells and increases the size of the Batchier network to accommodate the recirculating cells.

As discussed before, output buffering is the most preferable approach. However, Banyan networks cannot directly implement it since at most one cell per cell time is delivered to every output. Possible ways to work around this problem include:

- increasing the speed of internal links
- routing groups of links together
- using multiple Banyan planes in parallel
- using multiple banyan planes in tandem or adding extra switching stages

4.7.2.4.2.2 Multiple-Path MINs

Apart from banyan networks, many types of MINs with multiple paths between inputs and outputs exist. Classical examples include the non-blocking Benes and Close networks, the cascaded banyan networks, and the randomized route banyan network with

load distribution (which eliminates internal buffering). Combining a number of banyan planes in parallel can also be used to form multipath MINs.

The multipath MINs achieve more uniform traffic distribution to minimize internal conflicts, and exhibit fault tolerance. However if cells can take independent paths with varying delays, a mechanism is needed to preserve the sequential ordering of cells of the same virtual connection at the output. Since this might involve considerable processing, it is better to select the path during connection setup and fix it during the connection. Special attention must be paid during path selection to prevent unnecessary blocking of subsequent calls. Widjaja and Leon-Garcia have proposed a helical switch-using multipath MINs with efficient buffer sharing. They have used a virtual helix architecture to force cells to proceed in sequence.

4.8 Switch Design Principles

From the preceding section, it can be seen that each design alternative has its own merits, drawbacks, and considerations. The general design principles and issues exposed in the last section are analyzed in more detail here.

4.8.1 Internal Blocking

A fabric is said to be internally blocking if a set of N cells addressed to N different outputs can cause conflicts within the fabric. Internal blocking can reduce the maximum possible throughput. Banyan networks are blocking, while TDM buses where the bus operates at least N times faster than the port speed are internally nonblocking. By the same concept, shared memory switches which can read and write at the rate of NV cells per second are internally non-blocking, since if N cells arrive for N different outputs, no conflicts will occur. Hence, to prevent internal blocking, shared resources must operate at some factor greater than the port speed. Applying this to banyan networks, the internal links need to run square root of N times faster than the highest speed incoming link. This factor limits the scalability and throughput of the switch. Coppo et al. has developed a mathematical model for analyzing the optimal blocking probability versus complexity tradeoff.

4.8.2 Buffering Approaches

Buffering is necessary in all design approaches. For instance, in a banyan network, if two cells addressed to the same output successfully reach the last switching stage at the same time, output contention occurs and must be resolved by employing buffering. The location and size of buffers are important issues that must be decided.

There are four basic approaches to the placement of buffers. These basic approaches are illustrated in **figure 4.6**. The literature abounds with comparative studies of these, augmented with numerous queuing analyses and simulation results. Uniform random traffic, as well as bursty traffic has been examined. Although each approach has its own merits and draw-backs, output queuing is the preferred technique so far.

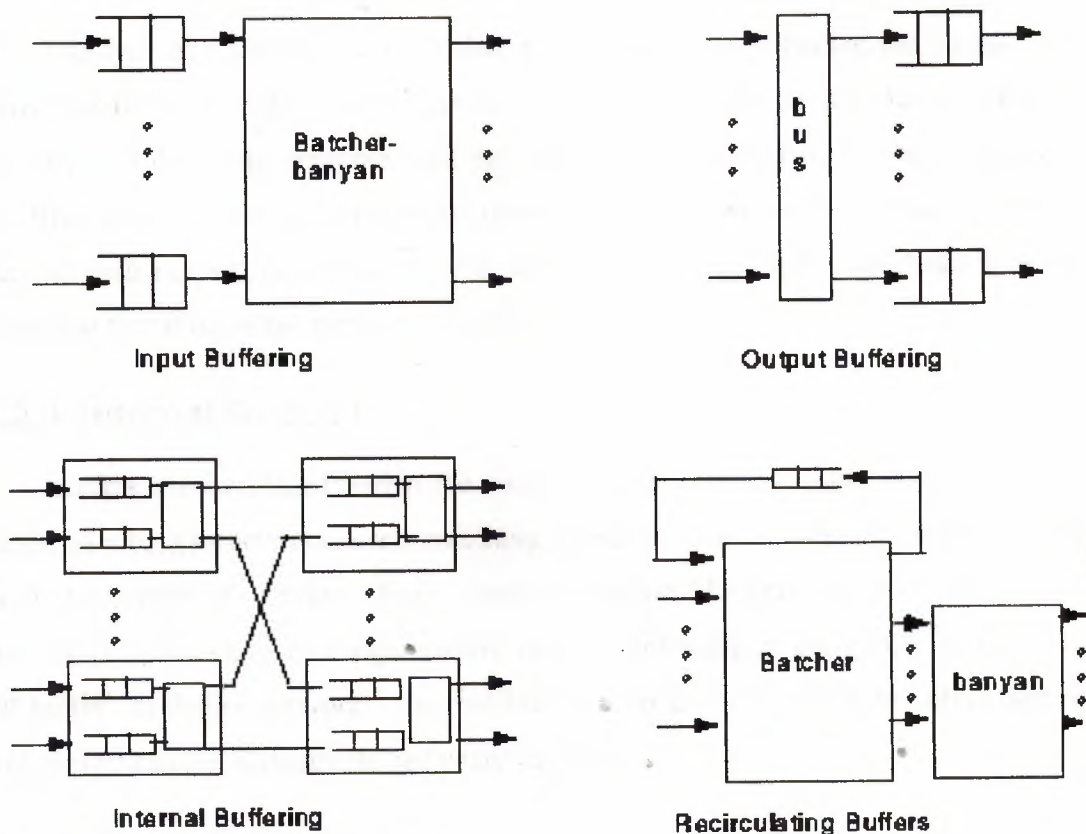


Figure 4.6: The various buffering approaches

4.8.2.1 Input Queuing

Buffers at the input of an internally nonblocking space division fabric (such as Batcher Banyan network) illustrate this type of buffering. This approach suffers from head-of-the-line blocking. When two cells arrive at the same time and are destined to the same output, one of them must wait in the input buffers, preventing the cells behind it from being admitted. Thus capacity is wasted. Several methods have been proposed to tackle the head-of-the-line blocking problem, but they all exhibit complex design. Increasing the internal speed of the space division fabric by a factor of four, or changing the First-In-First-Out (FIFO) discipline are two examples of such methods.

4.8.2.2 Output Queuing

This type of buffering can be evident by examining the buffers at the output ports of a shared bus fabric. This approach is optimal in terms of throughput and delays, but it needs some means of delivering multiple cells per cell time to any output. Hence, either the output buffers must operate at some factor times the port speed, or there should be multiple buffers at each output. In both cases, the throughput and scalability are limited, either by the speedup factor or by the number of buffers.

4.8.2.3 Internal Queuing

Buffers can be placed within the switching elements in a space division fabric. For instance, in a banyan network, each switching element contains buffers at its inputs to store cells in the event of conflict. Again, head-of-the-line blocking might occur within the switching elements, and this significantly reduces throughput, especially in the case of small buffers or larger networks. Internal buffers also introduce random delays within the switch fabric, causing undesirable cell delay variation.

4.8.2.4 Recirculating Buffers

This technique allows cells to re-enter the internally nonblocking space division network. This is needed when more than one cell is addressed to the same output simultaneously, so the extra cells need to be routed to the inputs of the network through the

recirculating buffers. Although this approach has the potential for achieving the optimal Throughput and delay performance of output queuing, its implementation suffers from two major complexities. First, the switching network must be large enough to accommodate the re-circulating cells. Second, a control mechanism is essential to sequentially order the cells.

4.8.3 Buffer Sharing

The number and size of buffers has a significant impact on switch design. In shared memory switches, the central buffer can take full advantage of statistical sharing, thus absorbing large traffic bursts to any output by giving it as much as is available of the shared buffer space. Hence, it requires the least total amount of buffering. For a random and uniform traffic and large values of N , a buffer space of only $12 N$ cells is required to achieve a cell loss rate of $1/10$ to the power of 9, under a load of 0.9.

For a TDM bus fabric with N output buffers, and under the same traffic assumptions as before, the required buffer space is about $90 N$ cells. Also a large traffic burst to one output cannot be absorbed by the other output buffers, although each output buffer can statistically multiplex the traffic from the N inputs. Thus buffering assumes that it is improbable that many input cells will be directed simultaneously to the same output. Neither statistical multiplexing between outputs or at any output can be employed with fully interconnected fabrics with N squared output buffers. Buffer space grows exponentially in this case.

4.8.4 Scalability of Switch Fabrics

If ATM switches will ever replace today's large switching systems, then an ATM switch would require a throughput of almost 1 Tbps. The problem of achieving such high throughput rates is not a trivial one. In all four switching design techniques previously analyzed, it is technologically infeasible to realize high throughputs. The memory access time limits the throughput attained by the shared memory and shared medium approaches, and the design exhibits a tradeoff between number of ports and port speed. The fully interconnected approach can attain high port speeds, but it is constrained by the limitations on the number of buffers.

The space division approach, although unconstrained by memory access time or number of buffers, also suffers from its own limitations:

1. Batcher-banyan networks of significant size are physically limited by the possible circuit density and number of input/output pins of the integrated circuit. To interconnect several boards, interconnection complexity and power dissipation place a constraint on the number of boards that can be interconnected
2. The entire set of N cells must be synchronized at every stage
3. Large sizes increases the difficulty of reliability and repairability
4. All modifications to maximize the throughput of space-division networks increase the implementation complexity

Thus the previous discussion illustrates the infeasibility of realizing large ATM switches with high throughputs by scaling up a certain fabric design. Large fabrics can only be attained by interconnecting small switch modules (of any approach) of limited throughput.

There are several ways to interconnect switch modules. The most popular one is the multistage interconnection of network modules. **Figure 4.7** illustrates the 3 stage Clos (N, n, m) network, which is a famous example of such an interconnection, and is used in Fujitsu's FETEX-150, NEC's ATOM and Hitachi's hyper-distributed switching system. There are N/n switch modules in the first stage, and each is of size $n \times m$. Thus the second stage contains m modules each of size $N/n \times N/n$, and the last stage again has N/n modules of size $n \times m$. Since this configuration provides m distinct paths between any pair of input and output, the traffic distribution can be balanced. Because each cell can take an independent path, cell sequence must be recovered at the outputs. Usually the least-congested path is selected during connection-setup. A new request cannot be accepted if the network is internally congested. Clos networks are strictly non-blocking if there always exists an available path between any free input-output pair, regardless of the other connections in the network. Since in ATM, the bandwidth used by a connection may change at different times, defining the nonblocking condition is not a trivial issue.

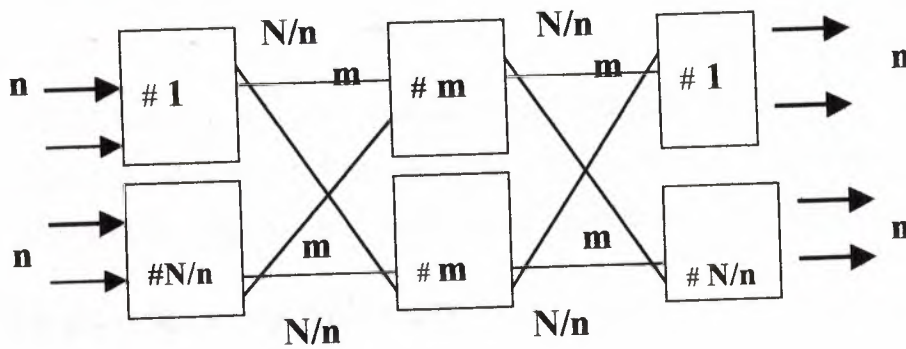


Figure 4.7: the clos (N, n, m) network

The throughput of the Clos network can be increased if the internal links have a higher speed than the ports. In that case, special care must be taken in selecting the size of the buffers at the last switching stage where most of the queuing occurs. The output buffering principle for Clos networks has been proposed to optimize through-put, where all buffering is located in the last stage of the interconnection network. This places an emphasis on the proper choice of the parameter m . m is usually selected according to the knockout principle previously discussed, which states that for a sufficiently large value of m , it is unlikely for more than m cells to arrive simultaneously for the same last stage module.

Another completely different approach for interconnection is to attempt to find the optimal partitioning of a large $N \times N$ fabric into small modules. The set of N inputs can be divided into K subsets, each handled by an $N/K \times N$ switch module. The K outputs are then multiplexed. The small switching modules in this case can be implemented as a Batcher sorting network, an expansion network or parallel banyan planes.

4.8.5 Multicasting

Many services, such as video, will need to multicast an input cell to a number of selected outputs, or broadcast it to all outputs. Designing multicasting capability can be done by either adding a separate copy network to the routing fabric, or designing each intercon-nected switching module for multicasting. Multicasting in different switch design is discussed next.

4.8.5.1 Shared Medium and Fully Interconnected Output-Buffered Approaches

Here multicasting is inherently natural, since input cells are broadcast anyway, and the address filters at the output buffers select the appropriate cells. Thus the address filters can simply filter according to the multicast addresses, in addition to the output port address.

4.8.5.2 Shared Memory Approach

Multicasting requires additional circuitry in this case. The cell to be multicast can be duplicated before the memory, or read several times from the memory. Duplicating cells requires more memory, while reading a cell several times from the same memory location requires the control circuitry to keep the cell in memory until it has been read to all output ports in the multicast group.

4.8.5.3 Space Division Fabrics

4.8.5.3.1 Crossbar Switches

Here multicasting is simple to implement but significantly impacts the switch performance. In crossbar switches with input buffering, broadcasting an input cell to multiple output ports is straightforward, but increases the head-of-the-line blocking at the input buffers. The only solutions to reduce this head-of-the-line blocking significantly increase the complexity of buffer control.

4.8.5.3.2 Broadcast Banyan Networks

In buffered banyan networks, multicasting can be realized if each switching element can broadcast an input cell to both outputs, while buffering another incoming cell. This technique, called a broadcast Banyan network (BBN), results in a number of complications. First, each switching element will have four possible states, and hence each cell will need two bits of control information at each stage of the Banyan. Furthermore, the two duplicate cells generated by the switching element are identical and thus would be routed to the same output port. This problem can be solved in either one of two ways. The first method would be to give multicast cells a multicast address instead of the output port address. This multi-

cast address can be transparently carried and used to determine routing information at each switching element. This method has the drawback of requiring more memory at each switching element. The second alternative would be to carry the entire set of output addresses in each multicast cell, so that the switching element can use this information to route and duplicate cells. This method obviously suffers from many problems.

4.8.5.3.3 Copy Networks

From the preceding discussion, it is clear that multicasting significantly increases the complexity of the space division fabric, and this has led several researchers to propose the separation of the cell duplication and routing functions into a distinct copy network and routing network. In this case, the copy network precedes the routing network. The copy network can recognize multicast cells and make the denoted number of duplicates, without being concerned with the careful routing of the duplicated cells. The routing network can then simply perform point-to-point routing, including routing all copies of the same cell (which is somewhat a drawback). But if the copy network follows the routing network, it will need to duplicate the cells, as well as deliver them, and we are back to the complexities of the broadcast network.

Implementing the copy network itself has been usually done by Banyan networks. The multicast cell is either randomly routed or duplicated at each switching stage, but duplication is delayed as much as possible to minimize resource usage. All duplicates of a cell will have the same addressing information, so the copy network will randomly route the cells. After the copy network, translation tables alter the multicast address to the appropriate output addresses. After observing that the broadcast Banyan network is nonblocking if:

- (a) the active inputs are concentrated,
- (b) there are N or fewer outputs
- (c) the sets of outputs corresponding to inputs are disjoint and sequential; concentration was proposed to be added before a broadcast Banyan networks.

The nonblocking feature of broadcast bananas also holds if the consecutiveness of inputs and outputs is regarded in a cyclic fashion. Byun and Lee investigate an ATM multicast

switch for broadcast Banyan networks, aimed at solving input fairness problems, while Chen et al. conclude that a neural network method for resolving output port conflict is better than a cyclic priority input access method for multi-cast switches.

4.8.6 Fault Tolerance

Because reliability is essential in switching systems, redundancy of the crucial components must be employed. The routing and buffering fabric, one of the most important elements of a switching system, can either be duplicated or redundant, and fault detection and recovery techniques must be implemented. Traffic can be distributed among the parallel fabric planes into disjoint subsets, partially overlapping subsets or it can be duplicated into identical sets. While the first approach provides the least redundancy, each plane carries only a small fraction of the total traffic, so it can be small and efficient. In contrast, duplicating the traffic into identical sets provides the greatest fault tolerance, but the least throughput. Partially overlapping subsets are a compromise.

Composing the routing and buffering fabric of parallel planes enhances fault tolerance, but adding redundancy within each individual fabric plane is still important. Designing fault-tolerant MINs has been the subject of a great deal of study. Banyan networks are especially fault-prone since there is only a single path between each input-output pair. Multipath MINs are more fault tolerant. To add redundancy to bananas, extra switching elements, extra stages, redundant links, alternate links or more input and output ports can be added. These techniques increase fault tolerance by providing more than one path between each input and output, and throughput of the Banyan can also increase. Of course, this is at the cost of a more complex implementation and routing. Time redundancy, where a cell attempts multiple passes through the MIN, has also been proposed. A Baseline tree approach has also been investigated by Li and Weng. This approach preserves most of the advantages of MINs, while providing a good performance under high traffic loads in the presence of faults in the switch.

A testing mechanism now needs to be implemented to verify the MIN operation and detect faults. A possible method can be to periodically inject test cells in a predetermined pattern and observe these cells at the outputs to detect the cell pattern. Addition of house-keeping

information to the cell header can also aid in discovering cell loss, misrouting or delays. Once a fault has been detected, traffic should be redistributed until the fault is repaired. The redistribution function can be carried out by concentrators, or by the MIN itself.

4.8.7 Using Priorities in Buffer Management

The cell switch fabric needs to handle the different classes of ATM traffic differently according to the Quality of Service (QOS) requirements of each. The traffic classes are mainly distinguished by their cell delay and cell loss priorities. Since output queuing was discovered to be the preferred approach, the switch fabric will have multiple buffers at each output port, and one buffer for each QOS traffic class. Each buffer is FIFO to preserve the order of cells within each VPC/VCC, but the queuing discipline must not necessarily be FIFO.

Buffer management refers to the discarding policy for the input of cells into the buffers and the scheduling policy for the output of cells from the buffers. These functions are a component of the traffic control functions handled by the Switch Management. Inside the switch fabric, the queues need to be monitored for signs of congestion to alert the Switch Management and attempt to control the congestion. A performance analysis of various buffer management schemes can be found in the paper by Huang and Wu. They also propose a priority management scheme to provide real-time service and reduce the required buffer size.

4.8.7.1 Preliminaries

The Cell Loss Priority (CLP) bit in the ATM cell header is used to indicate the relative discard eligibility of the cells. When buffers overflow, queued cells with CLP=1 are discarded before cells with CLP=0. The CLP bit is either set by the user to denote relatively low priority information, or set by the UPC when the user exceeds the traffic agreed upon.

Several degrees of delay priority can be associated with each virtual circuit connection. Since this is not part of the ATM cell header, it is usually associated with each VPI/VCI in the translation table within the switch. It can also be part of the internal routing tag

associated with each cell within the switch. Of course, cells of the same VCC must have the same delay priority, though they can have different cell loss priorities.

4.8.7.2 Cell Scheduling

Cell scheduling sets the order of transmission of cells out of the buffers. Since various QOS classes will usually have varying cell delay requirements, higher priority needs to be given to the class with more strict constraints. Static priorities whereby the lower priority class will be output only when there is no higher priority traffic has proved to be an unfair, as well as inflexible approach. This is because a large burst of higher priority traffic can cause excessively long queuing delays for the lower priority traffic.

A better approach is the deadline scheduling, where by each cell has a target departure time from the queue based on its QOS requirements. Cells missing their deadlines might be discarded based on switch implementation and traffic requirements. If service is given to the cell with most imminent deadline, the number of discarded cells can be minimized. A different scheduling scheme divides time into cycles, and takes scheduling decisions only at the start of each cycle, instead of before each cell.

4.8.7.3 Cell Discarding

Because cells of the same VPC/VCC must be maintained in sequential order, cells of different cell loss priorities might be mixed in the same buffer. A policy is needed to determine how cells with $CLP=0$ and cells with $CLP=1$ are admitted into a full buffer. In the push-out scheme, cells with $CLP=1$ are not admitted into a full buffer, while those with $CLP=0$ are admitted only if some space can be freed by discarding a cell with $CLP=1$. This scheme has an optimal performance. In contrast, the partial buffer-sharing approach entails that the cells with $CLP=0$ and $CLP=1$ can both be admitted when the queue is below a given threshold. When the queue exceeds this threshold, only cells with $CLP=0$ can be admitted as long as there is buffer space available. This can lead to inefficiency since $CLP=1$ cells can be blocked even when buffer space is available. This scheme can be designed for a good performance, and its implementation is much simpler than the push-out

one. Choudhury and Hahne propose a technique for buffer management that combines the backpressure mechanism with a push-out scheme to achieve a very low cell loss rate.

4.8.7.4 Congestion Indications

Buffer management will monitor queue statistics and alert the Switch Management if congestion is detected within the fabric. Queue statistics must be indicative enough for the Switch Management to determine whether congestion is increasing or receding, and whether it is focused or widespread. Buffer management must thus examine several fields in the internal routing tag, such as timestamps and housekeeping information.

Buffer management should be able to provide the Switch Management with performance data, congestion information, records of discarded cells, and usage management data for accounting. When congestion is detected, the Switch Management may instruct the buffer management to adjust the cell scheduling and discarding policies. The Explicit Forward Congestion Indication (EFCI) can be triggered, and in this case buffer management must alter the payload type field in the ATM cell header. Explicit rate schemes can also be activated for better congestion control. Most currently available switches only provide EFCI congestion control, such as ForeSystems ForeRunner and ForeThought ASX families.

CHAPTER 5

ATM Technology Components: Upper layers

Overview

So far we have discussed the transport and switching of cells. This chapter will go on with more describing and discussing of the behavior and functions of the ATM Adaptation Layer, the highest layer of the ATM protocol stack in specific and its types and services. And continue with the LAN emulation.

5.1 ATM Adaptation Layer Functions

The ATM adaptive layer (AAL) provides the interface from the user to the ATM system. As stated earlier, different data types require different types of AAL's. This is due to the characteristics of the data. Real time data such as voice and high resolution video can handle some loss of data but only low delay and bounded jitter delay. Non real time data, however, can handle hefty delays but cannot handle any losses. **Figure 5.1** shows the data rate and duration of various types of transmission. These new types of communication media impose new types of traffic loads on the networks and network components. The adaptability of ATM to handle all types of data is what makes it so applicable to today's needs.

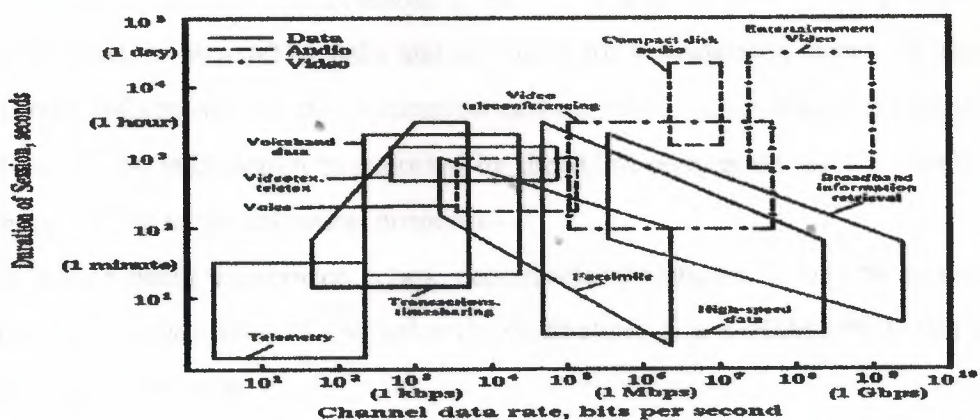


Figure 5.1. Data Rate and Duration of Various Transmissions

the desired QOS. According to a recent ATM Forum newsletter, work is currently being done on the development of a sixth AAL for MPEG2 video streams.

	AAL1	AAL2	AAL3	AAL4	AAL5
Timing relation between source and destination	required	required	not required	not required	not required
Bit rate	constant	variable	variable	variable	variable
Connection mode	connection oriented	connection oriented	connection oriented	connection-less	connection oriented

Figure 5.2 ATM Adaptation Layers

Along with providing the interface for the user, the AAL provides timing recovery, synchronization, and indication of loss of information, as well as other functions associated with a specific AAL. The actual connection as well as the connection characteristics is established by the AAL management functions.

The ATM adaptation layer lies between the ATM layer and the higher layers which use the ATM service. Its main purpose is to resolve any disparity between a service required by the user and services available at the ATM layer. The ATM adaptation layer maps user information into ATM cells and accounts for transmission errors. It also may transport timing information so the destination can regenerate time dependent signals. As mentioned earlier the information transported by the ATM adaptation layer is divided into four classes according to the following properties-

- 1) The information being transported is time dependent/independent: It may be necessary to regenerate the time dependency of a signal at the destination. E.g. a 64kbps PCM voice.
- 2) Variable/Constant bit rate.
- 3) Connection/Connectionless mode information transfer.

These properties define eight possible classes, four of which are defined as B-ISDN service

classes. Four ATM adaptation layer services are defined to match up with the four B-ISDN information classes: The ATM adaptation layer is divided into two sublayers:

1) Convergence Sublayer:

This layer wraps the user-service data units in a header and trailer which contain information used to provide the services required. The information in the header and trailer depends on the class of information to be transported but will usually contain error handling and data priority preservation information.

2) Segmentation and reassembly sublayer:

This layer receives the convergence sublayer protocol data unit and divides it up into pieces, which it can place in an ATM cell. It adds to each piece a header which contains information used to reassemble the pieces at the destination.

5.2 ATM Adaptation Layer Types

The types are originally classified according to the service to support:

- AAL 1: for constant bit-rate service
- AAL 2: for variable bit-rate service
- AAL 3: for Connection oriented service
- AAL 4: for connectionless data service

They are now classified according to the way the adaptation is performed and the need.

- AAL 1: is still supported
- AAL 2: not interesting and no progress on standardization
- AAL 3/4: AAL 3 and 4 are so similar in the format that they are merged
- AAL 5: A better way than AAL 3/4 for data communication
- SAAL: For Signaling purpose

5.2.1 AAL Type 1 Services And Functions.

AAL-1 (class A) is for synchronous bit streams; it is used to transfer constant bit rate data, which is time dependent. It must therefore send timing information between source and destination with the data so that the time dependency maybe recovered. AAL-1

provides error recovery and indicates errored information which could not be recovered. AAL-1 performs some functions such as: SAR, CS, handling cell delay variation, handling lost and misinserted cells, handling bit errors.

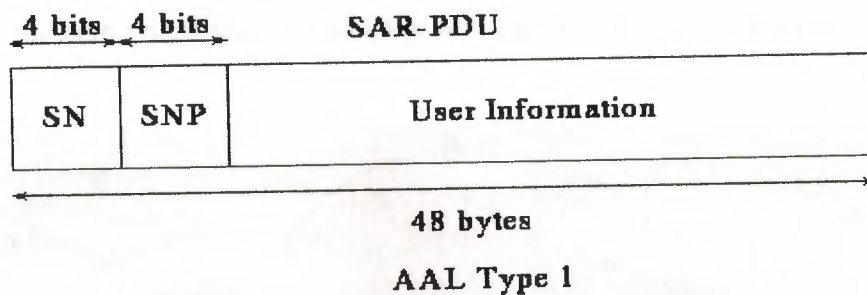
Convergence Sublayer:

The functions provided at this layer differ depending on the service provided. It provides bit error correction and may use explicit time stamps to transfer timing information.

- Data is broken into 47-byte cells
- A block is 124 cells + 4 error correcting cells (Reed-Solomon code)
- Each cell is given to the SAR sublayer to transmit.

Segmentation and reassembly sublayer:

At this layer the convergence sublayer protocol data unit is segmented and a header added. The header contains 3 fields



SN = Sequence Number

SNP = Sequence-Number Protection

SAR = Segmentation-And-Reassembly sublayer.

PDU = Protocol Data Units.

Figure 5.3 AAL-1 PDU field

- Sequence Number used to detect cell insertion and cell loss.
- Sequence Number protection used to correct and errors that occur in the sequence number.
- Convergence sublayer indication used to indicate the presence of the convergence sublayer function.
- The first byte of the 48 bytes data is used for SAR purpose
- The bits in the first byte:

- 1 bit: Convergence sublayer indicator (CSI): start of a block
- 3 bits: Sequence Count (SC): for detecting lost cells
- 3 bits: CRC for the first nibble: for detecting error in the SC
- 1 bit: even parity bit for the previous 7 bits: same as above

The receiver in ALL-1 has two modes of operation: correction mode and detection mode as in **fig. (5.4)**. The receiver starts out in correction mode, which is the default mode. In correction mode, the ALL-1 receiver is capable of providing single-bit error correction on the SAR-PDU header (not the SAR-PDU payload). If no error detected in correction mode, meaning that there is a valid SN based on the CRC and parity bit, no action is taken. If a single-bit error is detected and corrected (using the CRC) or a multiple-bit error is detected (same method), the receiver transitions to the detection mode, even if the corrected cell is still accepted and used by the receiver as a now valid cell.

In detection mode, all subsequent errored SNs are discard to prevent "strings" of cells needing constant processing. A valid SN will transition the receiver back to correction mode.

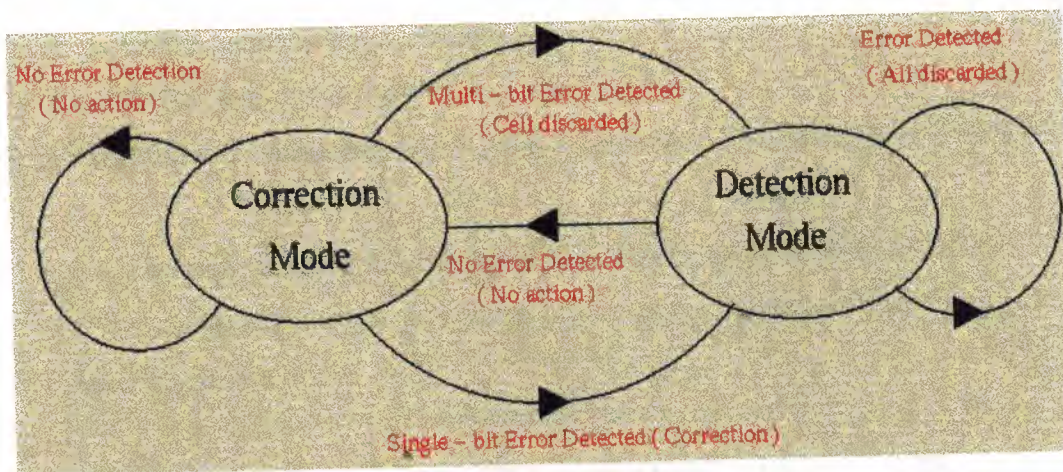


Figure 5.4 SNP receiver mode

5.2.2 AAL Type 2 Services And Functions.

This AAL (class B) was designed to support services such as compressed audio and video. AAL-2 is used to transfer variable bit rate data, which is time dependent. It sends timing information along with the data so that the timing dependency may be recovered at

the destination. AAL-2 provides error recovery and indicates errored information, which could not be recovered. As the source generates a variable bit rate some of technical cells transferred maybe unfull and therefore additional features are required at the segmentation and recovery layer.

Convergence sublayer:

This layer provides for error correction and transports the timing information from source to destination. This is achieved by inserting time stamps or timing information into the convergence sublayer protocol data unit.

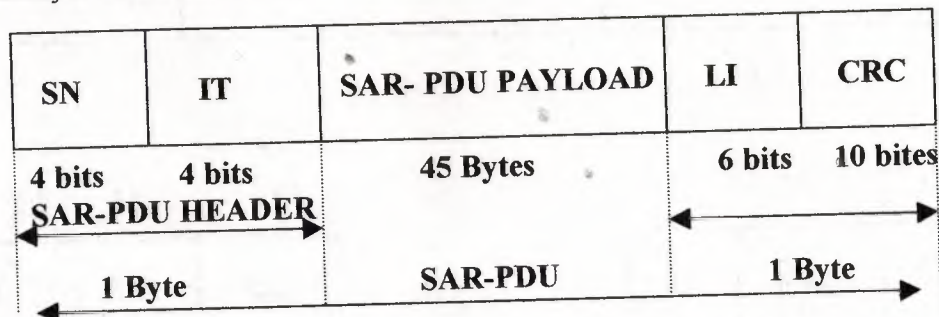
Segmentation and recovery sublayer:

The CS-PDU is segmented at this layer and a header and trailer added to each piece. The header contains two fields.

- 45 bytes of payload per cell
- Sequence number is used to detect inserted or lost cells.
- Information type is one of the following:
 - BOM , beginning of message
 - COM , continuation of message
 - EOM , end of message

or indicates that the cell contains timing or other information. The trailer also contains two fields:

- length indicator indicated the number of true data bytes in a partially full cell.
- CRC is a cyclic redundancy check used by the segmentation and reassembly sublayer to correct errors.



CRC Cyclic Redundancy Check

IT Information Type

LI Length Indicator

PDU Protocol Data Unit

SAR Segmentation And Reassembly

SN Sequence Number

Figure 5.5 ALL-2 PDU field.

5.2.3 AAL Type 3 Services And Functions.

AAL3 and AAL4 were originally conceived as separate protocols, but were later combined to form AAL3/4. AAL3 (Class C) is designed to transfer variable rate data, which is time independent. It supports both message mode and streaming mode services. Message mode services are transported in a single ATM adaptation layer interface data unit, while streaming mode services require one or more AAL-IDUs. AAL-3 can be further divided into two modes of operation:

- 1) Assured operation: Corrupted or lost convergence sublayer protocol data units are retransmitted and flow control is supported.
- 2) Non-assured operation: Error recovery is left to higher layers and flow control is optional.

Convergence sublayer:

The ATM adaptation layer 3 convergence sublayer is similar to the ATM adaptation layer 2 convergence sublayer as both handle non-real time data. The ATM adaptation layer 3 convergence sublayer is therefore subdivided into two sections:

- 1) The common part convergence sublayer. This is also provided by the AAL-2 CS. It appends a header and trailer to the common part convergence sublayer protocol data unit payload (as shown in **fig.5.6**)

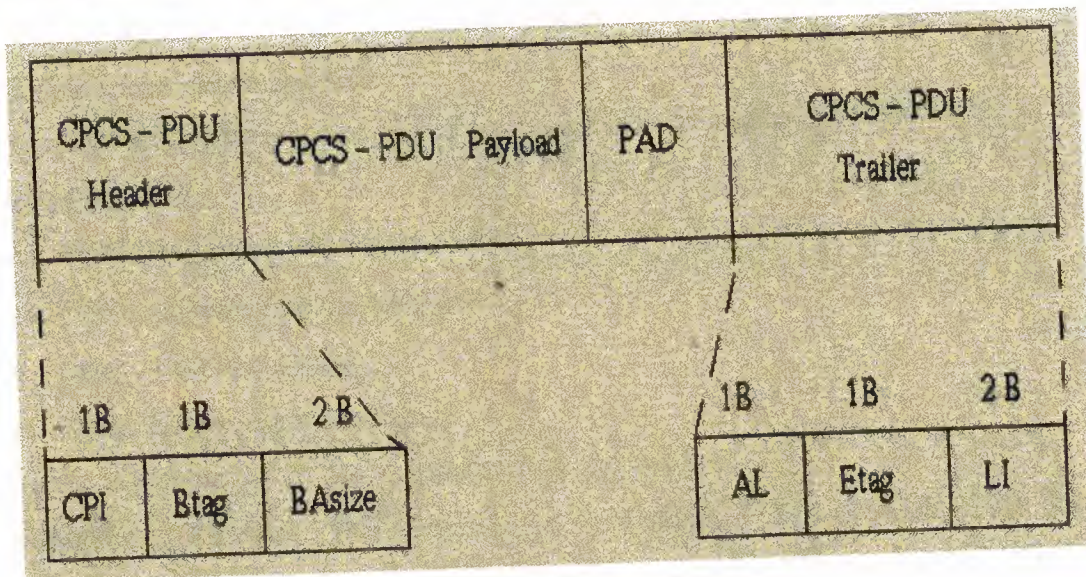


Figure 5.6.1

The header contains 3 fields:

- Common part indicator indicates that the payload is part of the common part.
- Begin tag marks the start of the common part convergence sublayer protocol data unit.
- Buffer allocation size tells the receiver how much buffer space is required to accommodate the message.

The trailer also contains 3 fields:

- Alignment is a byte filler used to make the header and trailer the same length. i.e. 4 bytes.
- End tag marks the end of the common part convergence sublayer - protocol data unit.
- The length field holds the length of the common part convergence sublayer protocol data unit payload.

2) The service specific part. The functions provided at this layer depend on the services requested. They generally include functions for error detection and recovery and may also include special functions such as transparent delivery.

Segmentation and reassembly sublayer: At this layer the convergence sublayer protocol data unit is segmented into pieces which can be placed in ATM cells. A header and trailer who contain information necessary for reassembly and error recovery are appended to each piece. The header contains 3 fields:

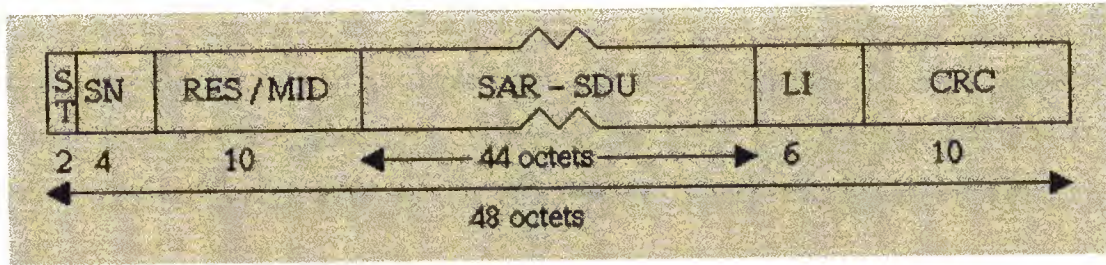
1) Segment type indicates what part of a message is contained in the payload. It has one of the following values:

- BOM : Beginning of message
- COM : continuation of message
- EOM : End of message
- SSM : Single segment message

2) Sequence number use to detect cell insertion and cell loss.

3) Multiplexing identifier. This field is used to distinguish between data from different connections, which have been multiplexed onto a single ATM connection. The trailer contains two fields:

- 1) Length indicator holds the number of useful data bytes in a partially full cell.
- 2) CYC is a cyclic redundancy check used for error detection and recovery.



ST Segment Type

SN Sequence Number

RES Reserved

Figure 5.6.2 AAL-3/4 PDU field.

5.2.4 AAL Type 4 Services And Functions.

AAL-4 (Class D) is designed to transport variable bit rate-time independent traffic in a connectionless mode. Like AAL-3 it supports to modes of service Message mode services and Streaming mode services. It also can operate in assured and non-assured mode (see ALL-3). AAL-4 provides the capability of transferring data with out establishing an explicit connection. It provides both point-to-point and point to multi-point transfer. The AAL cannot support a full connectionless service since functions like routing and network addressing are performed at a higher level.

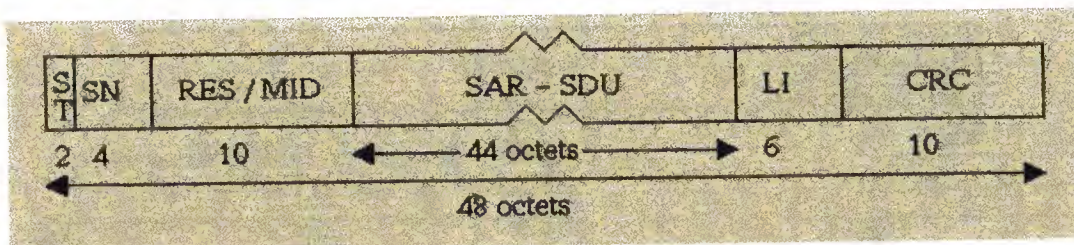


Figure 5.6.3 AAL-5 PDU Field

5.2.5 AAL Type 5 Services And Functions.

ALL-5 (Additional Class C) some times referred to as simple and efficient adaptation AAL5 is used to support Class C variable-bit-rate connection oriented data when high throughput is of greater concern than an occasional retransmission. AAL5 is

streamlined to minimize overhead. The error detection and multiplexing facilities of AAL3/4 are dropped to provide more room in each cell for the actual data payload.

These services provided by the ATM Adaptation Layer are just building blocks that allow various applications to use the underlying capabilities of ATM. In many applications, software is needed to implement a portion of the AAL. The long-term prospect is that hardware devices will largely absorb this function, increasing speed and driving costs down.

The Convergence Sublayer (CS)

- Data is padded and appended with an 8-byte trailer (total=48*N):
 - 1 byte: User-to-User Indication: Currently, this field is 0.
 - 1 byte: Common Part Indication: Currently, this field is 0.
 - 2 byte: Length: The number of bytes in the original packet
 - 4 byte: CRC-32 for the entire convergence sublayer packet

Segmentation and reassembly (SAR)

- Packets are broken into 48-byte units and put into the cells.
- The Payload Type field in the last cell's header is 1.

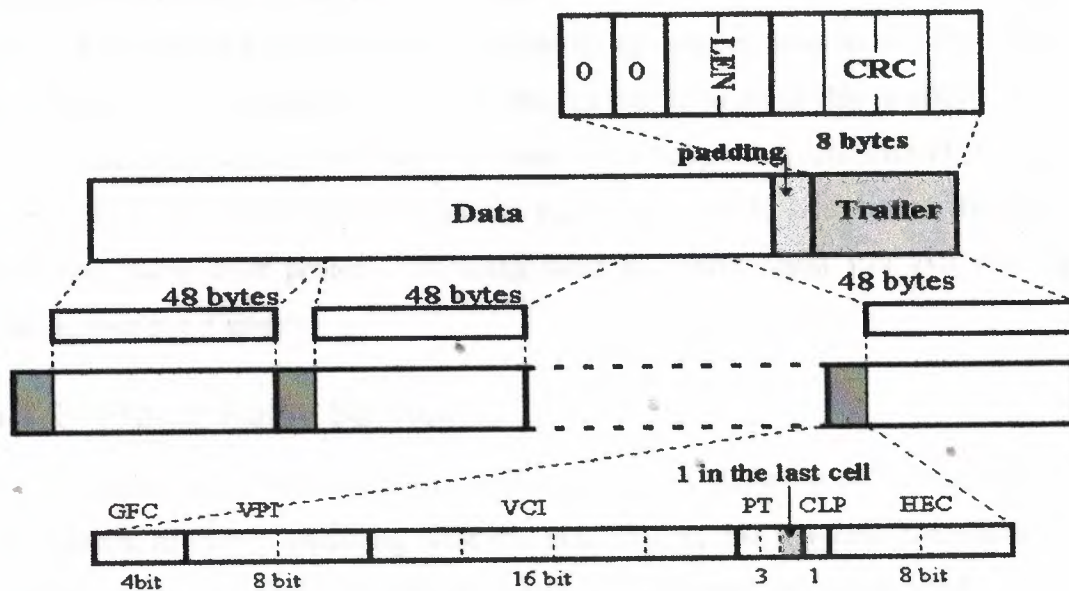


Figure 5.7 AAL-5 PDU fields

5.3 ATM Services

ATM network provides another types of services despite of AAL class services. That's one reason of the existing of ATM which to provide services to end-users. And here the end users do not care about classes or AALs or even cells. Users care about end to end: What can we use network for? . Many data services are presented to make the way to communi-cate with the users easy. Cell relay-relay service will just connect "ATM island" based on private "cell switched" LAN hubs and/or possibly routers. Frame-relay services will offer a cell-based backbone while giving users a familiar variable-length packet interface into the network. Connectionless network access protocol (CLNAP), switched multimegabit net-work services (SMDS), and connectionless broadband data service (CBDS) are three con-nectionless protocols for LAN interconnectivity and other connectionless services.

5.3.1 Cell-Relay Service (CRS)

Cell relay services (CRS) is an important initial offering of public ATM network providers and private ATM network builders. The cells already exist-in the ATM hub or router or multiplexer or even an ATM board in a work station on the user's premises. The public ATM network services offered would be to connect these local ATM "island" over larger areas than the private organization may be willing to go or able to afford.

Cell-relay services will not need any of the supporting structure of the AAL itself. CRS will be connection-oriented. All the sender need to do is to take a 48-byte payload from any higher-layer protocol providing them and add a valid VPI/VCI value in a cell header, then out it goes.

5.3.2 Frame Relay Service

Frame relay is a service that combines, in a way similar to ATM, the high speed and low latency of circuit switching with the port sharing and dynamic bandwidth allocation capabilities of traditional X.25 packet switching. Both frame relay and ATM are connection-oriented services that are using virtual channels to transport information. Today frame-relay is a widely used service for the bursty data traffic that characterizes most data applica-

tions, such as LAN interconnection. Frame relay can be introduced into multi-service networks where ATM provides the high-speed backbone. Network providers can optimize their networks by offering frame relay service on ATM switches. The end-user can use frame relay for lower-speed access and ATM for higher speeds and for multimedia applications. Interworking between frame relay and ATM is done via the ATM switch. There are two ways of providing frame relay services over ATM:

- Network interworking:** Network interworking uses the ATM backbone to interconnect two frame relay users. An interworking function (IWF) provides all the mapping and capsulation functions that are required. Each frame relay virtual circuit can be carried over an ATM virtual circuit or all of the frame relay virtual circuits can be multiplexed onto a single ATM virtual circuit.

Traffic enters the ATM network over a frame relay user network interface. The ATM network maps frame relay virtual circuits into ATM virtual circuits, and segments frame relay frames into ATM cells. The cells are transported through the ATM network to the destination node where the cells are reassembled into frames and handed over to the users via the frame relay user network interface.
- Service interworking:** Service interworking allows frame relay end-users to communicate with ATM end-users across the ATM network. The interworking function (IWF) performs the translation from ATM to frame relay and vice versa. With service interworking the traffic enters the network via the frame relay user network interface and exits over an ATM user network interface or vice versa. Transparent interoperability between frame relay and ATM users is provided. The interworking function (IWF) translates frame relay and ATM specific parameters that are used for traffic management and for quality of service. The way of encapsulation of higher level protocols (e.g. IP) is being converted by the IWF. Implementing frame relay / ATM network and service interworking on ATM switches offers the user a seamless migration path from frame relay to ATM. The move from frame relay to ATM can be done in a cost-effective way without affecting the service. A user that has migrated to ATM can still communicate with those users that are still connected by frame relay.

5.3.3 Connectionless Services (SMDS)

SMDS (Switched Multimegabit Data Service) has been specified as a data service for public carriers in the metropolitan area. SMDS is a connectionless high-speed data transmission service for the interconnection of LANs through public networks. The connectionless aspect of SMDS is its major benefit. Because SMDS is a connectionless service, users need not plan for and establish specific circuits between communicating parties for which maximum information rates would have to be defined.

Any user on an SMDS network can address a message to any other user, providing that both users tell their SMDS carriers to pass messages from their proposed partner. SMDS filtering provides free interconnection among partners without the risk of injection of unauthorized messages.

SMDS service can be delivered on top of an ATM infrastructure by providing SMDS interfaces plus the connectionless server function on ATM switches. To setup connectionless traffic over the connection-oriented ATM infrastructure requires at least one connectionless server or more depending on the traffic volume. The connectionless server function (CLS) provides the implementation of SMDS service over ATM networks and the interworking with existing SMDS networks. Scalability is achieved by distributing as many CLS in the backbone network as required by the end-user traffic.

This concept allows the flexibility that can grow smoothly with the connectionless traffic on the network. Figure 6.5 shows the concept of providing SMDS service over ATM. For the installed base of SMDS networks the connectionless server concept provides the evolution to ATM. New carriers can build multiservice networks by using an ATM platform offering SMDS service as part of the whole range of services that ATM has to offer.

5.3.4 LAN Emulation Services (LES)

5.3.4.1 The Need For LAN Emulation.

In order to use the vast base of existing LAN application software, it is necessary to define an ATM service, herein called "LAN Emulation", that emulates services of existing LANs across an ATM network and can be supported via a software layer in end systems.

Indeed, if such a LAN Emulation service is provided for an ATM network, then end systems (e.g. workstations, servers, bridges, etc.) can connect to the ATM network while the software applications interact as if they are attached to a traditional LAN. Also, this service will support interconnection of ATM networks with traditional LANs by means of today's bridging methods. This will allow interoperability between software applications residing on ATM-attached end systems and on traditional LAN end systems. The LAN Emulation service will be important to the acceptance of ATM, since it provides a simple and easy means for running existing LAN applications in the ATM environment.

5.3.4.2 LAN-Specific Characteristics To Be Emulated

5.3.4.2.1 Connectionless Services.

LAN stations today are able to send data without previously establishing connections. LAN Emulation provides the appearance of such a connectionless service to the participating end systems.

5.3.4.2.2 Multicast Services.

The LAN emulation service supports the use of multicast MAC addresses (e.g. broadcast, group, or functional MAC addresses). The need for a multicast service for LAN Emulation comes from classical LANs where end stations share the same media. Note, that supporting broadcast/multicast traffic does not necessarily mean that all messages addressed to a multicast MAC address must be distributed to every station. A large number of today's LAN protocols use broadcast or multicast messages. A service could be established to intercept these messages and forward them directly to their destinations instead of broadcasting them to every station. A simpler alternative would be to forward multicast messages to all stations and then rely upon filtering in those stations, as is done in existing LANs. This simpler approach is adopted in LAN Emulation.

5.3.4.2.3 MAC Driver Interfaces In ATM Stations.

The main objective of the LAN emulation service is to enable existing applications to access an ATM network via protocol stacks like APPN, NetBIOS, IPX, AppleTalk etc.

as if they were running over traditional LANs. Since in today's implementations these protocol stacks are communicating with a MAC driver, the LAN emulation service has to offer the same MAC driver service primitives, thus keeping the upper protocol layers unchanged. There are today some "standardized" interfaces for MAC device drivers: e.g. NDIS (Network Driver Interface Specification), ODI (Open Data-Link Interface) and DLPI (Data Link Provider Interface), which specify how to access a MAC driver. Each of them has its own primitives and parameter sets, but the essential services/ functions are the same. LAN Emulation provides these interfaces and services to the upper layers.

5.3.4.2.4 Emulated LANs.

In some environments there might be a need to configure multiple, separate domains within a single network. This requirement leads to the definition of an "emulated LAN" which comprises a group of ATM-attached devices. This group of devices would be logically analogous to a group of LAN stations attached to an Ethernet/IEEE 802.3 or 802.5 LAN segment. Several emulated LANs (ELANs) could be configured within an ATM network, and membership in an emulated LAN is independent of where system is an end physically connected. An end system could belong to multiple emulated LANs. Since multiple emulated LANs over a single ATM network are logically independent a broadcast frame originating from a member of a particular emulated LAN is distributed only to the members of that emulated LAN.

5.3.4.2.5 Interconnection with existing LANs.

As mentioned before, the LAN emulation service provides not only connectivity between ATM-attached end systems, but also connectivity with LAN-attached stations. This includes connectivity both from ATM stations to LAN stations as well as LAN stations to LAN stations across ATM. MAC layer LAN Emulation is defined in such a way that existing bridging methods can be employed, as they are defined today. Note, that bridging methods include both Transparent Bridging and Source Routing Bridging.

5.3.4.3 Description Of LAN Emulation Service

5.3.4.3.1 Basic Concepts.

LAN Emulation enables the implementation of emulated LANs over an ATM network. An emulated LAN provides communication of user data frames among all its users, similar to a physical LAN. One or more emulated LANs could run on the same ATM network. However, each of the emulated LANs is independent of the others and users cannot communicate directly across emulated LAN boundaries. Note that communication between emulated LANs is possible only through routers or bridges (possibly implemented in the same end station). Each emulated LAN is composed of a set of LAN Emulation Clients (LE Clients, or LECs) and a single LAN Emulation Service (LE Service). Each emulated LAN is one of two types:

- Ethernet/IEEE 802.3
- Token Ring/IEEE 802.5

LAN Emulation may be used in one of the two configurations:

- End Stations (e.g. PC or workstations)
- Intermediate Systems (e.g. Bridges or LAN Switches)

5.3.4.3.2 LAN Emulation In End Stations.

When a standard Ethernet or Token Ring application on an end station wishes to send data frame to the network via the standard software interface, the frame contains destination MAC address which uniquely identifies the destination, and this information is sufficient for the network adapter to transmit the frame. When this standard application operates over the ATM network the following steps should be performed:

- For a specified MAC address the network adapter must determine whether a VCC to the destination in the ATM network is already established. For this purpose a table of mappings between MAC addresses and VCCs is maintained.
- If there is no VCC established then first the ATM address of the destination (which is different from the MAC address) is obtained via the address resolution process, and then a VCC to the destination is set

up by the ATM network. This VCC is

known as "Data Direct VCC". The mapping's table is then updated.

- Once the station has established the Direct Data VCC the LANE header is appended to the frame, the frame is segmented into the 48-byte sized sells which are transmitted along the VCC.
- At the destination the sell stream is re-assembled and the original frame is recreated. The frame is then passed to the Ethernet or Token Ring application as if it was received through the standard interface.

5.3.4.3.3 LAN Emulation In bridges And LAN Switches.

LAN Emulation is used in the intermediate systems like Bridges and LAN Switches to enable physical Ethernet or Token Ring segments to interconnect with each other and with end stations across the ATM network. These devices can be thought of as a special kind of an end station which represents a number of MAC addresses - the MAC addresses of the stations attached to the LAN segment.

The Bridges and LAN Switches conceptually perform the actions of transferring the frames from one segment to the other according to the frame destination MAC address or the route information in the frame. When a physical LAN segment is to be connected to the ATM emulated LAN the Bridge with two inter-faces -ATM emulation and Standard LAN Interface -is required. The Bridge then receives the frames from the physical LAN and applies standard logic to decide whether to forward the frame. If a frame is to be forwarded, then the mechanism of MAC address to VCC mapping is used and the further frame handling is identical to that described above for end stations.

5.3.4.3.4 The Address Resolution Process.

One of the important processes is the process by which the end station or bridge resolves a LAN MAC address to an ATM address. For this purpose we need one other process, which is known as the LAN Emulation Server (LES), and protocol for communicating with the LES, which is known as the LAN Emulation Address Resolution Protocol (LE_ARP). We identify any end station or bridge that implements LAN Emulation as a LAN emulation Client (LEC), so the LEC is a process that stay in the end station or

bridge, which provides the entry point to emulated LAN. When the LAN Emulation Client want to know the ATM address of another LEC that it knows the MAC address for, it send a request to LES using LE_ARP protocol. If the LES knows what ATM address matches the requested MAC address, it send this information to the LEC. If it does know, it forwards the request to all other LECs.

Some of these LECs may be bridges, and they will respond to the E_ARP request if the MAC address matches one that they know about on their attached LAN segments. For this to work, it is necessary for each LEC to have VCC set up to the LES on which to send LE_ARP request, and on which to receive responses. This VCC, known as the "Control Direct VCC", is set up by each LEC when it joins the emulated LAN.

The join process takes place when the end station or bridge that implemented LAN emulation is started up. During the join process, a LEC will exchange information with LES so that the LES can maintain the table which consists of all the LECs currently active on the emulated LAN. The LES is the forwarding of LE_ARP requests to multiple LAN Emulation Clients when it is not able to resolve the request based on its own data. This distribution could be carried out via the individual control direct VCCs that the LES maintains with each LEC. But it's efficient to use point-to-multipoint VCC, and so the LES can be chosen to set up a point-to-multipoint VCC for all LECs as they join. This is known as the "Control Distribute VCC".

5.3.4.3.5 Forwarding Of Broadcast And Multicast Frames.

Let's discuss now the special case of broadcast and multicast frames. In fact, the process for dealing with this kind of frames in the end stations is just the same as for unicast frames: the MAC broadcast or multicast destination address is mapped to VCC, and the frame is segmented into 48 byte cell payloads for transmission on this VCC.

However, the ATM network now has to take care of ensuring that the broadcast or multicast frame reaches all its intended destinations. It does this by means of an additional process in the ATM network known as the Broadcast and Unknown Server (BUS). When a LEC registers with the LES, it immediately issues a LE_ARP request to find the ATM address that corresponds with the broadcast MAC address (hex FFFFFFFF). The LES responds to this LE_ARP request with the ATM address of the BUS.

The end station then proceeds to set up a VCC to the BUS. This is known as the "Multicast Send VCC". Once the LEC has established this connection to the BUS, the BUS registers a return path to the LEC, where possible by adding the LEC to an existing point-to-multipoint VCC. This is known as the "Multicast Forward VCC".

Alternatively, the BUS may establish a point-to-point VCC to the LEC. Whenever the application in the end station requests the transmission of a broadcast or multicast frame, the LEC process in that end station uses this VCC to send the frame to the BUS. When the BUS receives a broadcast or multicast frame from a LEC, it copies the frame and forwards it to all the LECs that are registered with it, using the Multicast Forward VCC(s). It should be noted that the BUS sends the broadcast and multicast frames to all the LECs that are registered with it, including the LEC that originates the broadcast. It is the responsibility of each LEC to discard any superfluous broadcast & quota; echoes". It can do this by looking for a match between the LAN Emulation header on the frame and its own unique LEC identity (LECID).

5.3.4.3.6 The LAN Emulation Configuration Server.

There may be many instances of the LAN Emulation Server (LES) in an ATM network, each one representing a single emulated Ethernet or Token Ring LAN segment. But how the LAN Emulation Client (LEC) finds the ATM address of the LES so that it can setup the Control Direct VCC and execute the registration process. One possible way of handling this would be to configure the address of the appropriate LES into each of the LECs. However this scheme is inefficient and inflexible, and so LAN Emulation provides us with the option of defining a LAN Emulation Configuration Server (LECS) which can be queried by any LEC to obtain the ATM address of the LES it can register with. A LEC wishing to locate a LES with which it can register sets up a Configuration Direct VCC with the LECS, specifies the range of parameters that may include the type of LAN emulated, the maximum packet size, the emulated MAC address of the LEC, and the name of emulated LAN which the LEC wishes to join. The LECS uses this information to determine which LES would be appropriate for this LEC, and responds with the ATM address of the relevant LES.

5.3.4.3.7 Network Architectures With LAN Emulation.

LAN Emulation over ATM allows for the implementation of emulated Ethernet and Token Ring segments that far exceed the limitations of real LAN in terms of physical reach, bandwidth and number of stations. ATM is not limited on physical distance. Using leased lines or public cell-switching services, It should be possible to construct emulated LANs that span the globe with a single emulated segment. ATM almost not limited on bandwidth or carrying capacity. Data rate between end stations connected via point-to-point VCC can reach 622 Mbps. And there is no theoretical limits on the number of stations that may be connected to an emulated LAN.

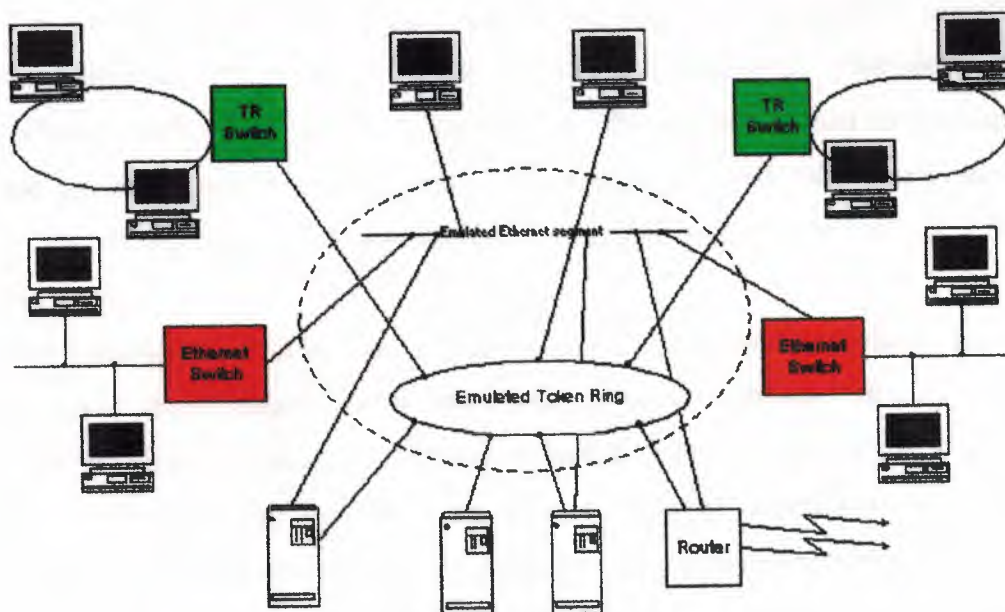


Figure 5.8 Logical views of integrated LAN - ATM environment

CHAPTER 6

ATM AND THE FUTURE

There has always been a general need for more and cheaper bandwidth, but nevermore so than today. This need has been accelerated by the availability of high-performance and relatively low-cost PCs and workstations and capable business software applications. In the future public network operators (PNOs) and service providers will be increasingly engaged in providing easy to manage broadband services on demand at an affordable price.

The situation now taking shape holds out new opportunities and the prospect of more widely differentiated roles for all new and traditional operators in the sector. The network offers a wide range of services, including information transport on stationary, mobile and satellite networks with national, international and global connectivity, supply of voice, data transmission, video distribution or integrated services and basic as well as value added services.

ATM development rests on a basis of well-established standards and on subsequent detail specifications. These specifications were formulated with the precise objective of providing detailed and widely accepted reference specifications to the international community of network and service operators, manufacturers, system integrators and users. They are used in developing interworkable commercial equipment capable of providing the advanced communication services required for the information society.

In July 1996, a major milestone in the development of ATM technology was reached when The ATM Forum announced agreement on the Anchorage Accord. The Anchorage Accord contains the Foundation Specs for mission-critical ATM infrastructure, as well as expanded Feature Specs for migration to ATM multi-service networks. In the future, foundation Specs will be revised to correct problems, align with ITU-Recommendations, and add cohesiveness/parity between specifications. Any new-specifications will be backward compatible Anchorage Accord specifications. The introduction of ATM-layer services will increase the benefits of ATM, making the technology suitable for a virtually unlimited range of applications, since an ATM

network is able to offer different levels of service. The corporate information technology and communication market has flexibility and a capacity for rapid change that are far superior to those of conventional telecommunications. The emerging ability to work together in an open environment (i.e. between solutions provided by different manufacturers) over a wide geographical area will increase the weight of this market.

In the corporate setting one of the priority requirements is the interconnection of Local Area Networks, or LANs. Offering this service on a public or private wide area network allows high-volume data and image transfer, e.g. for computer-aided design and manufacturing applications, workgroup sessions, or multimedia communications, where correlated processing of data, image and voice information is performed by the user terminal. In the area of corporate services the most pressing need is to be able to make good use of the speed, availability, quality and security which have become such marked features of the local area networks. As far as future home-applications are concerned, multimedia will be an essential feature and, particularly if associated with interactivity is destined to have a profound impact on our habits and social dealings.

This certainty aside, it is still too early to tell how quickly this market will widely develop, given its heavy dependence on mass phenomena. Nevertheless it is now time to begin assessing the appeal of new services and new ways of providing services on sample groups of customers. Given that the residential sector has not yet made much use of narrow band ISDN services, it is hard to see an immediate demand for residential broadband services. One exception is the home entertainment sector. The demand for improved quality cable television and simultaneous growth in the installation of residential fibre access offers a chance for integration, at least at the access level, of residential communication and entertainment needs. This will provide a starting point for future broadband services. Many studies based on questionnaires and analyses of cable TV services have shown that, in principle at least, subscribers are also willing to pay for various video retrieval services and enhanced quality TV.

There are several reasons why ATM is the basic technological choice for high-speed multi-service networks. This choice, endorsed internationally, is aimed at developing a general, unifying solution for all application environments. Here the primary objective is integrated handling of signals produced by voice sources, data,

images and video, regardless of the bandwidth associated with each type of information. In principle a network based on ATM can thus satisfy the requirements of both the business environment and the household environment. Demand for broadband services is in its initial phase and is growing. While just at the beginning of the commercial offer by network operator and service providers, inherent speed, response time and quality advantages will ensure the gradual uptake of services based on broadband networks.

Conclusion

In this project I have tried to analyze different issues pertaining to ATM, and discussed some of the ideas of different researches done on it. ATM is a relatively new field, and is still in taking its efficiency. Researching and building experimentally ATM networks, and implementing different protocols over it. ATM-Type networks play an important role in the broadband communications networks now and in the future.

This project has been compiled to make such an examination possible. It describes the wonderful resource and achievement that done on the field of networking and what it has become. It is hoped that this project will encourage many who are not yet online to realize that ATM networking is well worth the difficulty and trouble of studying and working on it. Also this project is intended to provide the material for courses for students in education, communication, library science, computer science and engineering, etc. about the exciting and important resources available to the society from this important technological breakthrough of ATM networking. The material contained in this project is intended to encourage students and teachers everywhere to begin to research and write about these developments so that there begins to be a substantial body of material documenting both where these developments have come from and the potential that they promise for a better world.

Most importantly, the information in this project are offered as a way of providing the public whose money and labor made these achievements possible, with a way of evaluating proposals to change the course of development for this network. These informations about ATM are a contribution toward evaluating what has been created, and what its social and scientific potential is. The past few years have seen people in Eastern Europe and around the world demonstrate that they need better living and working conditions. These cries for change mean that the methods that have achieved this Global network need to be applied to other aspects of society so that the well being of the people become the concern of government and of the public arena. Thus this collection of information about ATM is for those who want to see the coming millenium bring a better world for everyone.

APPENDIX I

AAL	ATM Adaptation Layer
ABR	Available Bit Rate
ACTS	Advanced Communications Technologies and Services
API	Application Program Interface
ARP	Address Resolution Protocol
ATM	Asynchronous Transfer Mode
B-ICI	B-ISDN Inter Carrier Interface
B-ISDN	Broadband-ISDN
BUS	Broadcast and Unknown Server
CATV	Cable TV
CBDS	Connectionless Broadband Data Service
CBR	Constant Bit Rate
CDV	Cell Delay Variation
CEC	Commission of the European Community
CES	Circuit Emulation Services
CLNAP	ConnectionLess Network Access Protocol
CLP	Cell Loss Priority
CLR	Cell Loss Ratio
CMIP	Common Management Interface Protocol
CNM	Customer Network Management
CS	Convergence sublayer
CSCW	Computer Supported Collaborative Work
CTD	Cell Transfer Delay
ECU	European Currency Unit
EEA	European Economic Area
ELAN	Emulated LAN
EMAC	European Market Awareness Committee
ETSI	European Telecommunications Standard Institute
EU	European Union

FDDI	Fibre Distributed Data Interface
FR	Frame Relay
FUNI	Frame based UNI
GCRA	Generic Cell Rate Algorithm
HEC	Internet Protocol
IPX	Internetwork Packet Exchange
ISDN	integrated services digital network
IT	Information Technology
ITU	International Telecommunication Union
LAN	Local Area Network
LANE	LAN Emulation
LEC	LAN Emulation Client
LECS	LAN Emulation Configuration Server
LES	LAN Emulation Server
LUNI	LAN Emulation UNI
MAC	Media Access Control
MAN	Metropolitan area networks
MBS	Maximum Burst Size
MCR	Minimum Cell Rate
MIB	Management Information Base
MPOA	MultiProtocol Over ATM
MSS	MAN Switching System
N-ISDN	Narrowband-ISDN
NNI	Network Node Interface
nrt-VBR	non real time VBR
OAM	Operation Administration and Maintenance
ONP	Open Network Provision
OSI	Open Systems Interconntection
PCR	Peak Cell Rate
PDH	Plesiochronous Digital Hierarchy
PICS	Protocol Implementation Conformance Statement

Appendix I

PM	Physical Medium
PNNI	Private NNI
PNO	Public Network Operator
PSTN	Public Switched Telephone Network
PT	Payload Type
PVC	Permanent Virtual Channel
QOS	Quality of Service
rt-VBR	real time VBR
SAR	Segmentation and Reassembly sublayer
SCR	Sustainable Cell Rate
SDH	Synchronous Digital Hierarchy
SEAL	Simple and Efficient Adaptation Layer
SMDS	Switched Multi-megabit Data Service
SME	Small and Medium Enterprises
STM	Synchronous Transfer Mode
TC	Transmission Convergence
TCP/IP	Transmission Control Protocol/Internet Protocol
TDM	Time Division Multiplexing
TIP	Transport and Interworking Package
TMN	Telecommunications Management Network
UBR	Unspecified Bit Rate
UNI	User Network Interface
UPC	Usage Parameter Control
VBR	Variable Bit Rate
VCC	Virtual Channel Connection
VCI	Virtual Channel Identifier
VCL	Virtual Channel Link
VLAN	Virtual LAN
VOD	Video-On-Demand
VP	Virtual Path
VPI	Virtual Path Identifier

Appendix I

VPN	Virtual Private Network
WAIS	Wide Area Information Service
WAN	Wide Area Network
WWW	World Wide Web

References

1. Black, Uyless.D. (1995). "ATM: Foundation for Broadband Networks", Prentice Hall, Inc.
2. Black, Uyless.D. "ATM Volume III/Internetworking With ATM".
3. Rainer. Handel, Huber. " ATM Networks Concepts Protocols Applications".
4. Onvural, Raif.O. (1994)."Asynchronous transfer mode networks: performance issues", Boston: Artech House.
5. Thomas M. Chen and Stephen S. Liu. (1995). "ATM Switching Systems", Artech House, Incorporated.
6. Dhas, Chris, Vijaya K. Konangi, and M. Sreetharan. (1991)." Broadband switching architectures, protocols, design, and analysis". Los Alamitos, Calif.: IEEE Computer Society Press.
7. [Ftp://ftp.netlab.ohio-state.edu/pub/jain/courses/cis788-95/atm-cog/index.html](http://ftp.netlab.ohio-state.edu/pub/jain/courses/cis788-95/atm-cog/index.html)
8. <http://ee.wpi.edu/courses/ee535hwkllcd95/bkh/refs.html>.
9. www.rad.com/networks/1194/atm.tutorial.htm