

NAVAL POSTGRADUATE SCHOOL

Monterey, California



THESIS

NETWORK POLICY MANAGEMENT

by

Paul A. Wetzel

September 1999

Thesis Advisor:
Associate Advisor:

Geoffery Xie
Dr. Gilbert Lundy

Approved for public release; distribution is unlimited.

DTIC QUALITY INSPECTED 3

20000502 066

REPORT DOCUMENTATION PAGE

Form Approved
OMB No. 0704-0188

Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instruction, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188) Washington DC 20503.

1. AGENCY USE ONLY (Leave blank)	2. REPORT DATE September 1999	3. REPORT TYPE AND DATES COVERED Master's Thesis	
4. TITLE AND SUBTITLE Network Policy Management		5. FUNDING NUMBERS	
6. AUTHOR(S) Paul A. Wetzel.			
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Postgraduate School Monterey, CA 93943-5000		8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES)		10. SPONSORING / MONITORING AGENCY REPORT NUMBER	
11. SUPPLEMENTARY NOTES The views expressed in this thesis are those of the author and do not reflect the official policy or position of the Department of Defense or the U.S. Government.			
12a. DISTRIBUTION / AVAILABILITY STATEMENT Approved for Public Release; Distribution is Unlimited.		12b. DISTRIBUTION CODE	
Effective military and commercial use of Internet to conduct mission critical and commerce oriented transactions over shared networks is increasingly inhibited by the shortcomings of the very enabling technology of the Internet - the TCP/IP protocol. Without network performance, security and other management controls, TCP/IP networks can't meet the overall requirements of a network. To complicate the network policy management issues, new applications are exchanging increasingly larger amounts of digital data (image, audio, video, etc.), and some of them require stringent Quality-Of-Service (QOS) measured by delay and loss from the network. This places very diverse but demanding requirements on the network in terms of bandwidth and data delivery dependencies on the network. In many cases this network traffic diversity has led to major network performance and reliability problems as well as a resultant loss of productivity among network users. [Ref. 1]. Typical 10Mbit Ethernet LANs or even 100Mbit switched LANs running TCP/IP are no longer adequate to handle the various types of next generation applications being written and the existing mission critical applications. The bottom line is that more than bandwidth is required and the existing network infrastructure installed base make it impossible to quickly change to a new standard such as ATM for all desktop connections. As a result there is a growing need for management tools capable of running a new generation of applications over existing infrastructure. These requirements have not gone unnoticed by Government and commercial enterprises, network infrastructure vendors or the standards bodies.			
14. SUBJECT TERMS Network Policy Management		15. NUMBER OF PAGES	
		16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT Unclassified	18. SECURITY CLASSIFICATION OF THIS PAGE Unclassified	19. SECURITY CLASSIFICATION OF ABSTRACT Unclassified	20. LIMITATION OF ABSTRACT UL

NSN 7540-01-280-5500

Standard Form 298 (Rev. 2-89)
Prescribed by ANSI Std. Z39-18

Approved for public release; distribution is unlimited

NETWORK POLICY MANAGEMENT

Paul A. Wetzel
Lieutenant Commander, United States Navy
B.S., Texas A&M University, 1986

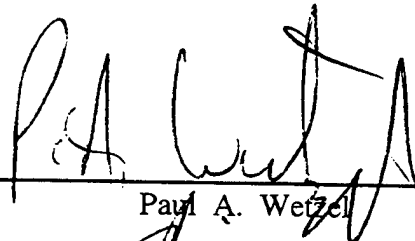
Submitted in partial fulfillment of the
requirements for the degree of

MASTER OF SCIENCE IN INFORMATION TECHNOLOGY MANAGEMENT

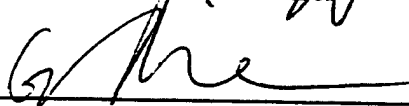
from the

**NAVAL POSTGRADUATE SCHOOL
September 1999**

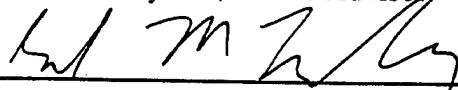
Author:


Paul A. Wetzel

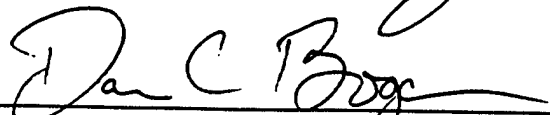
Approved by:



Dr. Geoffrey Xie, Thesis Advisor



Dr. Gilbert Lundy, Associate Advisor



Dr. Dan C. Boger, Chairman
Information Systems Academic Group

ABSTRACT

Effective military and commercial use of Internet to conduct mission critical and commerce oriented transactions over shared networks is increasingly inhibited by the shortcomings of the very enabling technology of the Internet - the TCP/IP protocol. Without network performance, security and other management controls, TCP/IP networks can't meet the overall requirements of a network. To complicate the network policy management issues, new applications are exchanging increasingly larger amounts of digital data (image, audio, video, etc.), and some of them require stringent Quality-Of-Service (QOS) measured by delay and loss from the network. This places very diverse but demanding requirements on the network in terms of bandwidth and data delivery dependencies on the network. In many cases this network traffic diversity has led to major network performance and reliability problems as well as a resultant loss of productivity among network users. [Ref. 1]

Typical 10Mbit Ethernet LANs or even 100Mbit switched LANs running TCP/IP are no longer adequate to handle the various types of next generation applications being written and the existing mission critical applications. The bottom line is that more than bandwidth is required and the existing network infrastructure installed base make it impossible to quickly change to a new standard such as ATM for all desktop connections. As a result there is a growing need for management tools capable of running a new generation of applications over existing infrastructure. These requirements have not gone unnoticed by

Government and commercial enterprises, network infrastructure vendors or the standards bodies.

TABLE OF CONTENTS

I.	INTRODUCTION	1
A.	PURPOSE	5
B.	RESEARCH QUESTIONS	6
C.	EXPECTED BENEFITS OF THIS THESIS	6
II.	BACKGROUND	9
A.	INTRODUCTION TO NETWORK MANAGEMENT	9
1.	Management Defined	9
2.	Network Management Defined and Discussed	9
a.	<i>Functional Architecture of Network Management</i>	13
(1)	Defining the Pieces	13
(2)	Managed Objects	14
(3)	Element Management Systems (EMS)	14
(4)	Manager of Managers Systems (MoM)	14
(5)	User Interface	15
b.	<i>Management Functional Areas (MFAs)</i>	15
(1)	Fault Management	16
(2)	Configuration Management	16
(3)	Accounting	17
(4)	Performance Management	17
(5)	Security	18
(6)	Chargeback	18
(7)	Systems Management	19
(8)	Cost Management	19
c.	<i>Common Implementations</i>	20
(1)	Management Focus	21
(2)	The Right Implementation	21
d.	<i>Business Case Requirements</i>	21
(1)	Definition	22
(2)	Levels of Activity	22
(3)	Today's Implementations	24
e.	<i>System Focus</i>	26
3.	Management Functional Domains (MFDs)	28
a.	<i>Introduction</i>	28
b.	<i>Building Requirements</i>	29
4.	A New Paradigm for Network Management	33
a.	<i>The Advent of Web-Based Management</i>	34
b.	<i>The Benefits of Web-Based Management</i>	35
c.	<i>Two Strategies for Implementation</i>	36
d.	<i>Emerging Standards</i>	38

	e.	<i>Web-Based Management and the World Wide Web</i>	40
5.		Questions to Ask	41
	a.	<i>How Much with the System Cost?</i>	41
	b.	<i>Will the Proposed System Integrate Into and Enhance My Current MIS Support Capabilities?</i>	41
	c.	<i>Is the Proposed System Modular in Design?</i>	42
	d.	<i>Is the Product Proposed Just an Element Management System or is it an Integrator of Element Management Systems?</i>	42
	e.	<i>What Does the System Monitor?</i>	42
	f.	<i>Does the Proposed System Enhance the Capabilities of the Current Support Staff or Does It Add More Support Staff?</i>	42
6.		Conclusion	43
B.		POLICY-BASED MANAGEMENT	43
	1.	Definition of a Network Policy	44
	2.	The Need for Policy Networking	49
	3.	Policy Networking Tree of Variables	51
III.		MANAGING QUALITY OF SERVICE	59
	A.	INTRODUCTION	59
	B.	DIVERGENT VIEWS	63
	C.	DEFINING QUALITY OF SERVICE	67
		1. Availability	69
		2. Performance	69
		3. Accuracy	70
		4. Affordability	70
	D.	APPROACHING THE QOS PROBLEM	71
	E.	SERVICE LEVEL AGREEMENTS	73
		1. Problems of the Past	74
		2. Value of SLAs	75
	F.	QOS SOLUTION	76
	G.	INTERNET PROTOCOL QOS MECHANISMS	78
		1. Integrated Services Model	79
		a. <i>Flows</i>	79
		b. <i>Service Categories</i>	79
		c. <i>Traffic Specification</i>	79
		d. <i>Requested Service Specification</i>	80
		e. <i>Path Characterization</i>	80
		f. <i>Resource Reservations</i>	81
		g. <i>Soft State</i>	82
		2. Delay and Jitter	84
		3. Throughput	85
		4. Issues with the int-serv Model	86
	H.	DIFFERENTIATED SERVICES MODEL EXPLORED	87

I.	INTERACTION WITH MPLS.....	92
J.	A NEW CHALLENGE: MANAGING QUALITY OF SERVICE (QOS)	94
1.	Queuing Techniques for Congestion Management on Outbound Traffic	95
a.	<i>First In, First Out (FIFO) Queuing: Basic Store and Forward</i>	96
(1)	Policy Requirements for FIFO Queuing Interfaces	97
(2)	FIFO's Relationship to Traffic Coloring	97
b.	<i>Priority Queuing (PQ): Basic Traffic Prioritization</i>	97
(1)	Policy Requirements for Priority Queuing Interfaces	98
(2)	Priority Queuing's Relationship to Traffic Coloring.....	98
c.	<i>Custom Queuing (CQ): Advance Traffic Prioritization</i>	98
(1)	Policy Requirements for Custom Queuing Interfaces	99
(2)	Custom Queuing's Relationship to Traffic Coloring.....	100
d.	<i>Weighted Fair Queuing (WFQ): Intelligent Traffic Prioritization</i>	100
(1)	Policy Requirements for Weighted Fair Queuing Interfaces	101
(2)	Weighted Fair Queuing's Relationship to Traffic Coloring	101
e.	<i>Weighted Round Robin (WRR): Traffic Taking Turns</i>	101
(1)	Modulo-N Hash	102
(2)	Hash-Threshold.....	102
(3)	Highest Random Weight (HRW)	102
2.	Traffic Shaping or Traffic Limiting Techniques for Controlling Bandwidth	103
a.	<i>Generic Traffic Shaping (GTS): Controlling Traffic on Non-Frame Relay Interfaces</i>	104
(1)	Interface QoS Property Requirements for Generic Traffic.....	104
b.	<i>Frame Relay Traffic Shaping (FRTS): Controlling Traffic on Frame Relay Interfaces and Subinterfaces</i>	105
(1)	Interface QoS Property Requirements for Frame- Relay Traffic Shaping	106
c.	<i>Limiting: Limiting Bandwidth and Optionally Coloring Traffic</i>	106
(1)	Interface QoS Property Requirements for Rate Limited Traffic.....	107
(2)	QoS Policy	107
IV.	NETWORK POLICY: SECURITY	109
A.	OVERVIEW OF NETWORK SECURITY.....	109

B.	A GROWING NEED FOR PROTECTION	110
C.	SECURITY STRATEGIES	112
	1. Least Privilege.....	112
	2. Defence in Depth.....	113
	3. Choke Point.....	114
	4. Fail Safe Stance	114
	5. Security Through Obscurity	115
	6. Simplicity	116
	7. Host Based Security	116
	8. Network Based Security	117
D.	SECURITY POLICY.....	117
	1. Site Security Policy.....	118
	2. Network Service Access Policy	119
	3. Firewall Design Policy	120
	4. System Specific Policies.....	121
	5. Incident Handling.....	122
	6. Disaster Recovery	122
E.	THE KEY ELEMENTS OF SECURITY.....	123
	1. Authenticity	123
	2. Access Control	124
	3. Integrity	124
	4. Confidentiality.....	124
	5. Security Must be Implemented in the Network.....	125
	6. Authentication Services	125
	a. Secure Access to Layer-Two Groups.....	127
	b. Device Authentication	129
	c. User Authentication	130
	(1) Authentication Server	131
	(2) Authentication Agent	133
	(3) Authentication Client	133
F.	FIREWALL SERVICES	134
	1. Types of Firewalls	135
	a. Packet Filtering Firewalls	135
	b. Application Gateways	135
	c. Circuit-level Gateways	136
	d. Stateful Inspection	136
	2. Switch-based Firewalls.....	137
	a. Port Count	137
	b. Performance	137
	c. Redundancy	138
	d. Management.....	138
	e. QoS	139
	3. External Firewalls	139
	4. Internal Firewalls	141
	5. Firewall Architecture.....	142
	6. Switch-embedded Firewall Features.....	143
	a. Access Control.....	144
	(1) Authenticity—Firewall Authentication	146
	(2) User Authentication	146
	(3) Client Authentication	146

	(4) Transparent Session Authentication	147
	b. Logging	147
	c. Address Translation	148
	(1) Dynamic Mode	149
	(2) Static Mode	149
G.	LOOKING AHEAD—DIRECTORY INTEGRATION	150
H.	TYPES OF COMPUTER SECURITY POLICY	150
	1. Program-level Policies	150
	2. Program-framework Policies	151
	3. Issue Specific Policies	151
	4. System-specific Policies	151
I.	EXAMPLES OF SECURITY POLICY	152
	1. Program-Level Policy	152
	a. <i>An Example of a Program Level Policy</i>	153
	2. Program-Framework Policy	154
	a. <i>Program-Framework Policy Example</i>	155
	3. Issue-Specific Policy	156
	a. <i>Issue-Specific Policy Example</i>	157
	4. System-Specific Policy	161
	a. <i>Security Objectives</i>	161
	b. <i>Operational Security</i>	161
	c. <i>Policy Implementation</i>	162
J.	THE SECURITY CHALLENGE	163
V.	NETWORK COST MANAGEMENT POLICY: (A CASE STUDY)	165
A.	EXECUTIVE SUMMARY	165
B.	NETWORK COST PROBLEM DEFINES	166
C.	BANDWIDTH ISSUES	167
D.	CONFIGURATION ISSUES	168
E.	DOWNTIME DILEMMA	169
F.	SECURITY ISSUES	170
G.	AN ORGANIZATION PLAN OF ACTION	170
	1. A Single-Site Solution	172
	2. The Multiple-Site Solution	182
H.	A PLAN TO FOLLOW	186
VI.	CONCLUSION: HIGH PERFORMANCE NETWORKING INITIATIVES	189
A.	NEXT GENERATION INTERNET NETWORK POLICY	
	APPLICATIONS AT NASA	189
	1. The Results of the Game Plan: NASA Participation at NGI	192
	a. Technologies	192
	b. Testbeds	193
	2. Naval Postgraduate School Contribution: Server and Agent	
	Based Active Management (SAAM) Architecture	196
	3. Other High-Performance Networking Initiatives	197
	3. Networking Policy Considerations for Next-Generation	
	Applications	199
	a. <i>Specific Policies</i>	201
	(1) QoS	201
	(2) Security	202

	(3) Policy Architecture.....	204
B.	RECOMMENDATIONS/SUGGESTIONS FOR FURTHER STUDIES AND RESEARCH TOPICS	207
	1. Gemini, an Example of the Need for Research Collaboration Linkages.....	207
	2. IPV6	209
	LIST OF REFERENCES	211
	BIBLIOGRAPHY	215
	INITIAL DISTRIBUTION LIST	225

LIST OF FIGURES

1.	Network Management Growth [Ref. 16]	12
2.	Levels Of Functionality [Ref. 18].....	13
3.	Network Management Center (each dot representing a MNC [Ref. 8] ..	20
4.	Distributed System [Ref. 8].....	25
5.	Distributed System [Ref. 8].....	27
6.	The Proxy Solution for WBM [Ref. 20].....	37
7.	The Embedded Approach to WBM [Ref. 20]	37
8.	A Procedure for Continuous Archival of Logging and Check-Pointing Data. [Ref. 17]	47
9.	Policy-Based Network Management Growth Policy Networking Building Blocks. [Ref. 12].....	50
10.	Network Policy Tree	52
11.	Security Policy Tree.....	52
12.	QoS Policy Tree	53
13.	QoS Integrated Services	54
14.	QoS Differential Services	55
15.	QoS Policy Mechanisms.....	56
16.	Cost Policy Tree	57
17.	Qos Architecture [Ref. 18]	77
18.	Policy Networking Architecture [Ref. 12]	95
19.	QoS Policy Binding [Ref. 21]	108
20.	Device Identification [Ref. 28].....	130
21.	User Identification [Ref. 28]	131
22.	External Firewall [Ref. 28]	140
23.	Internal Firewall [Ref. 28].....	142
24.	Firewall Architecture [Ref. 27]	143
25.	Corporate Enterprise 1994 [Ref. 34].....	172
26.	Corporate Enterprise 1997 [Ref. 34].....	178
27.	Bandwidth Graphs [Ref. 34]	179
28.	Support Staff Ratio [Ref. 34].....	183
29.	Power Network 1997 [Ref. 34].....	185
30.	Power Network 1998 [Ref. 34].....	186
31.	Access Levels of Hierarchy [Ref. 34].....	205

LIST OF TABLES

1	Availability. [Ref. 19].....	65
2.	A Switch-Based Firewall Has Two Primary Applications: as an External Firewall and as an Internal Firewall. [Ref. 30].....	140

ACKNOWLEDGMENT

The author would like to acknowledge those individuals who provided their support throughout the information gathering phase of this thesis....Thank-You Kiersten for being there when I needed you. Also thanks to Hank, Elvis, Virgil and Tulpen for expressing yourselves and your vocal support.

I. INTRODUCTION

The management of information technology projects has long been a significant problem for federal agencies. The government obligated more than \$23.5 billion toward information technology products and services in fiscal year 1994--about five percent of the government's total discretionary spending. Yet the impact of this spending on agency operations and service delivery has been mixed at best. Federal computer systems often cost millions more than expected, take longer to complete than anticipated, and fail to significantly improve the speed and quality of federal programs--or reduce their cost. Some private and public sector organizations, however, have significantly improved performance by managing their information technology resources within an overall framework that aligns technology with business needs and priorities [Ref. 1].

There is an increasing demand, coming from the Congress and the public, for a smaller government that works better and costs less. Having valuable, accurate, and accessible financial and programmatic information is a critical element for any improvement effort to succeed. Furthermore, ~~increasing~~ the quality and speed of service delivery while reducing costs will require the government to make significant investments in three fundamental assets-- personnel, knowledge, and capital property/fixed assets.

Investments in information technology (IT) projects can dramatically

affect all three of these assets. Indeed, the government's ability to improve performance and reduce costs in the information age will depend, to a large degree, on how well it selects and uses information systems investments to modernize its often-outdated operations. However, the impact of information technology is not necessarily dependent on the amount of money spent, but rather on how the investments are selected and managed. This, in essence, is the challenge facing federal executives: Increasing the return on money spent on IT projects by spending money wiser, not faster.

IT projects, however, are often poorly managed. For example, one market research group estimates that about a third of all U.S. IT projects are canceled, at an estimated cost in 1995 of over \$81 billion [Ref. 1]. In the last 12 years, the federal government has obligated at least \$200 billion for information management with mixed results at best. Yet despite this huge investment, government operations continue to be hampered by inaccurate data and inadequate systems.

Too often, IT projects cost much more and produce much less than what was originally envisioned. Even worse, often these systems do not significantly improve mission performance or they provide only a fraction of the expected benefits. Of 18 major federal agencies, 7 have an IT effort that has been identified as high risk by either the Office of Management and Budget (OMB) or the Congress [Ref. 2].

Some private and public sector organizations, on the other hand, have designed and managed IT to improve their organizational performance.

In a 1995 report, congress analyzed the information management practices of several leading private and state organizations [Ref. 3]. These leading organizations were identified as such by their peers and independent researchers because of their progress in managing information to improve service quality, reduce costs, and increase work force productivity and effectiveness. From this analysis, fundamental IT management practices were identified that, when taken together, provide the basis for the successful outcomes found in leading organizations.

One of the best practices exhibited by leading organizations was that they manage information systems projects with clear-cut management policies [Ref. 4].

This particular practice offers organizations great potential for gaining better control over their IT expenditures and performance. In the short term (within 2 years), this practice serves as a powerful tool for carefully managing and controlling IT expenditures and better understanding the explicit costs and projected returns for each IT project. In the long term (from 3 to 5 years), this practice serves as an effective process for linking IT projects to organizational goals and objectives. However, managing IT projects with specific policies works most effectively when implemented as part of an integrated set of management practices. For example, project management systems must also be in place, reengineering improvements analyzed, and planning processes linked to mission goals.

The Congress has passed several pieces of legislation that lay the

groundwork for agencies to establish a policy approach for managing IT. For instance, revisions to the Paperwork Reduction Act (PRA) (Public Law 104-13) have put more emphasis on evaluating the operational merits of information technology projects. The Chief Financial Officers (CFO) Act (Public Law 101-576) focuses on the need to significantly improve financial management and reporting practices of the federal government. Having accurate financial data is critical to establishing performance measures and assessing the returns on IT investments. Finally, the Government Performance and Results Act (GPRA) (Public Law 103-62) requires agencies to set results-oriented goals, measure performance, and report on their accomplishments.

In addition, the recently passed Information Technology Management Reform Act (ITMRA) (Division E of Public Law 104-106) requires federal agencies to focus more on the results achieved through IT investments while streamlining the federal IT procurement process. Specifically this act, which became effective August 8, 1998, introduces much more rigor and structure into how agencies approach the selection and management of IT projects. Among other things, the head of each agency is required to implement a process for maximizing the value and assessing and managing the risks of the agency's IT acquisitions.

Defense computing networks are not the only networks needing improvement in management.

Businesses have flocked to the Internet, deployed intranets at a dizzying rate and pushed out massive multimedia applications to desktop users. Amid this flurry of activity, many businesses failed

to anticipate the network congestion these new applications can cause [Ref. 5].

Typical 10Mbit Ethernet LANs or even 100Mbit switched LANs running TCP/IP are no longer adequate to handle the various types of next generation applications being written and the existing mission critical applications [Ref. 6]. The bottom line is that more than bandwidth is required and the existing network infrastructure installed base make it impossible to quickly change to a new standard such as ATM for all desktop connections. Instead of funnelling monetary resources into creating more bandwidth "pipe", network managers must find methodologies to optimize use of existing network resources, including those that have been plentiful and have been taken for granted [Ref. 7]. Policy based network management is an emergent management technology that meets this requirement-transmission prioritization. This thesis will examine network policy management and its uses.

A. PURPOSE

The purpose of this thesis is to: (1) provide a review of network policies in practice, (2) research how these network policies are enforced, (3) provide an evaluation of DOD network policies, (4) provide an evaluation of commercial network policies, (5) provide an overview of the structure of Management Information Systems (6) examine network resource management methodologies, (7) review the most current research being conducted by DoD and commercial entities, including NASA, in computer networking and (8) explore the benefits of this research.

B. RESEARCH QUESTIONS

- What are the historical/traditional economic decision variables/parameters used to define network policy management?
- What variables are required to manage current distributed computing networks?
- What is the effect of additional/differing decision variables on network policy management methodologies?
- What are the commonly practiced computing network policies of network administrators in academic institutions?
- What are the commonly practiced computing network policies of network administrators in military institutions?
- What are the commonly practiced computing network policies of network administrators in the business sector?
- How are the commonly practiced computing network policies enforced?
- How are commonly used Internet protocols related to network policy management and how are they categorized?
- What would be an appropriate way for DOD network managers to administer their bandwidth and prioritize network usage?

C. EXPECTED BENEFITS OF THIS THESIS

This research will assist network managers in assessing the impact of emerging technologies that utilize prioritization mechanisms. Readers will benefit from presentation of a clear definition of network resources, which is unavailable from other information sources. They will be provided a snapshot of prioritization methodologies that are in use and available. This research will explain the prioritization and policy options, their usage, and the advantages and disadvantages of each methodology. This thesis intends to expand on the traditional views of network management and will show how "policy

management" aims to reduce human intervention in networking. The reader will understand that by giving the network a set of rules ("policies") that govern how it will respond to certain events, then building on and maintaining that set of rules, we can replace the case-by-case manual configuration that takes place today. The research will also provide insight into alternatives to continued high-cost infrastructure upgrades within DOD computing networks.

THIS PAGE INTENTIONALLY LEFT BLANK

II. BACKGROUND

A. INTRODUCTION TO NETWORK MANAGEMENT

This chapter will provide a detailed background of network management and policy management basic concepts. This section (A) will define basic terminology, functional areas, and tasks included in a typical network management model. The term "management" is defined below to build a basis for further detailed discussion of specific network management techniques.

1. Management Defined

According to the American Heritage Dictionary (1997), to manage is:

- To direct or control the use of, to handle.
- To exert control over.
- To make submissive to one's authority, discipline, or persuasion.
- To succeed in accomplishing or achieving, especially with difficulty; contrive or arrange.

Managing IT resources is a complicated and expensive task, which includes many areas of responsibility. In the business sense, management has always connoted optimizing the utility of available resources - in other words - making the most of current assets. In effect, the IT manager's responsibility is to control the use of his or her resources to create the greatest profit for an organization.

2. Network Management Defined and Discussed

Network Management as a term has many definitions dependent on the operational functions in question. This chapter will illustrate and discuss today's

most common implementations of network management systems as they apply to actual MIS (Management Information System) form and function. It will also illustrate a *What's wrong with this picture* type of scenario, then discuss what the ideal system will look like.

Network management systems have been in operation many years especially in their own proprietary worlds such as Tivoli's Netview, AT&T Accumaster and Digital Equipment Corporation's DMA. With the implementation of Simple Network Management Protocol (SNMP), local area and wide area network components could be monitored and "managed". With the vast amount of raw data available, most MIS Managers have no idea what they really want because, in part, they don't know what is available [Ref. 10]. Additionally, how does the data get into a format that actually means something? Some communications systems are even considered non-manageable because they are only accessible by an RS-232 port and not by Netview or SNMP. There exists a belief that Network Management means nothing but the monitoring and management of network architectural hardware such as Routers, bridges and concentrators -- nothing above the network layer of the OSI (Open Systems Interconnection) model is considered manageable [Ref. 10].

It is alarming that most Senior Network Engineers tend to be resigned to spend thousands of dollars on hardware and software BEFORE the real requirements are gathered and defined. Consequently, MIS departments either spend very little on network management or they "go for broke" with the huge

hardware platforms and expensive artificial intelligence engines driving network management for the company [Ref. 10].

In today's environment of cost cutting and productivity enhancements, many common network management implementations call for an increase the number of people required to support MIS functions. These new people are senior level engineering and support types; very expensive in most cases. Typical costs extend into the hundreds of thousands of dollars purchasing hardware and software, not to mention the expense for additional personnel.

Network management systems have to be geared toward the workflow of the organization in which they will be utilized [Ref. 11]. As each MIS implementation is geared toward the business requirements, so should the network management system. If the management functionality does not directly or indirectly solve a business problem, it is totally useless to the overall MIS department and to the company.

Network management does not mean one application with a database with some huge *chunk of iron* running the show. It is really an integrated conglomeration of functions that may locate on one machine, or span thousands of miles, different support organizations and many machines and databases. Each of these functions must be directly driven by the business case. Most network decision makers agree that the real concern of the future is not necessarily migrating to high-speed intranets and extranets, but reducing the spiraling costs associated with supporting and operating the network [Ref. 11].

Figure 1 graphically depicts an estimation of the growth of network management market.

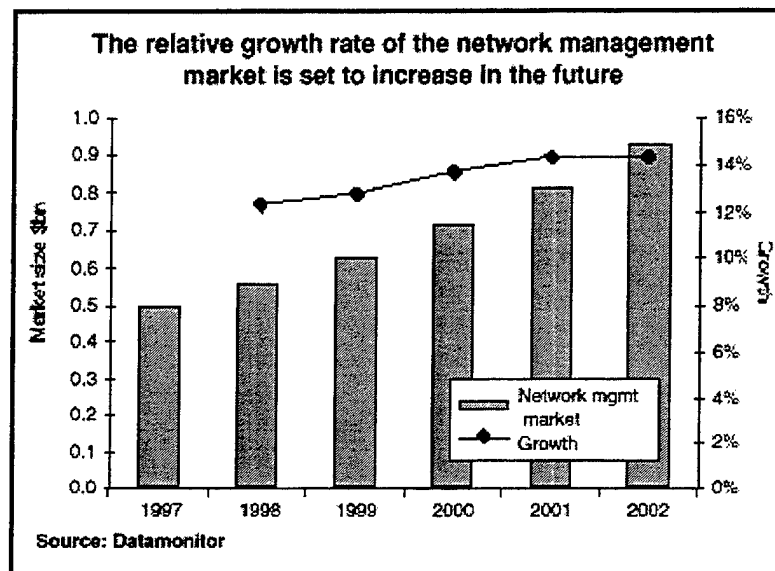


Figure 1. Network Management Growth. [Ref. 16]

a. Functional Architecture of Network Management

(1) Defining the Pieces. Network management systems have four basic levels of functionality. Each level has a set of tasks defined to provide, format, or collect data necessary to manage the objects. Figure 2 illustrates these four levels of functionality.

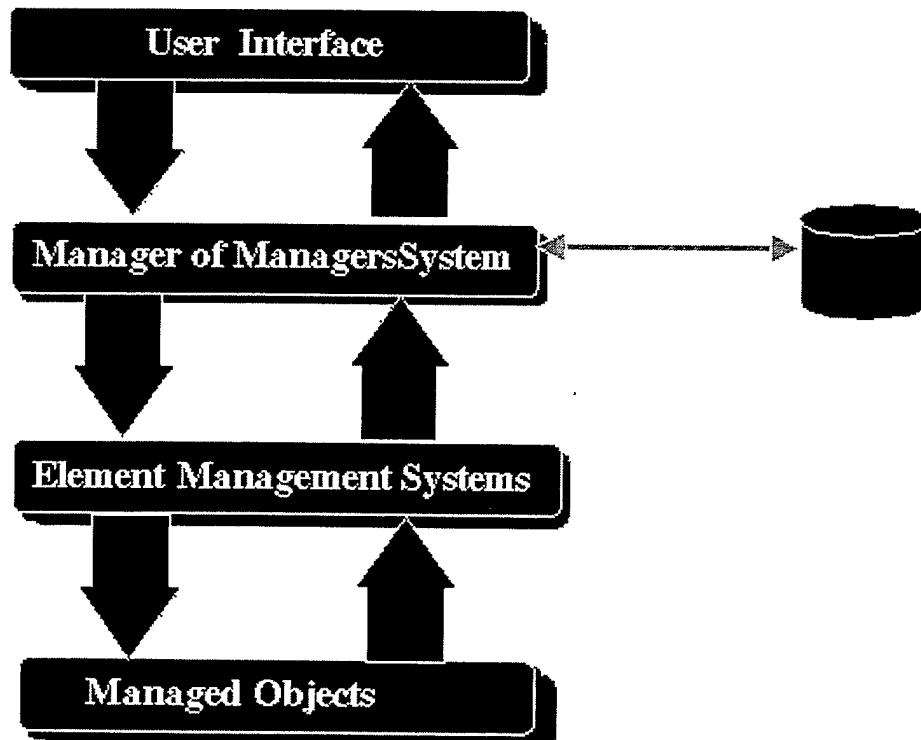


Figure 2. Levels of Functionality. [Ref. 18]

(2) Managed Objects. Managed Objects are the devices, systems and/or anything else that require some form of monitoring and management. Most implementations leave out the "anything else" clause because the implementors usually do not have the business case requirements before the design; therefore they design as they go.

Some examples of managed objects include routers, concentrators, hosts, servers and applications like Oracle, Microsoft Management Systems, Lotus Notes, and Microsoft Mail. The managed object does not have to be a piece of hardware but should rather be depicted as a function provided on the network [Ref. 10].

(3) Element Management Systems (EMS). An EMS manages a specific portion of the network. For example, the SunNet Manager, an SNMP management application, produced by Sun Microsystems, is often used to manage SNMP manageable elements. Element Managers may manage multiplexers, Private Automatic Branch eXchange (PABX's), proprietary systems or an application [Ref. 10].

(4) Manager of Managers Systems (MoM). MoM systems integrate together the information associated with several element management systems, usually performing alarm correlation between EMS's. There are several different products that fall into this category, including Boole & Babbage's CommandPost, NyNEX 's AllLink, International Telematic's MAXM, OSI's

NetExpert and others.

(5) User Interface. The user interface to the alarm information, whether real time alarms and alerts or trend analysis graphs and reports, is the principal piece to deploying a successful network management system. If the information gathered cannot be distributed to the whole MIS organization to keep people informed and to enable team communications, the real purpose of the system is lost in the implementation [Ref. 10]. Data doesn't mean anything if it is not used to make informed decisions about the optimization of network resource usage.

The above systems components are, in turn, mapped back to what is called Management Functional Areas (MFAs). These MFAs are the wish list of areas in which management applications, as a system should focus.

b. Management Functional Areas (MFAs)

The most common framework depicted in network management designs is centered around the OSI "FCAPS" model of MFAs. Most network management implementations do not cover all of these areas. Other less inclusive models may not address all areas that may be important to the MIS function and to specific business units within the company [Ref. 10].

FCAPS is an acronym explained as follows:

Fault Management

Configuration Management

Accounting

Performance Management

Security Management

Some of the other areas covered under the FCAPS model are:

Chargeback

Systems Management

Cost Management

(1) Fault Management. Fault management is the detection of a problem, fault isolation and correction to normal operation. Most systems poll the managed objects, search for error conditions and illustrate the problem in either a graphic format or a textual message. Most of these types of messages are setup by the person configuring the polling on Element Management Systems [Ref. 10]. Some Element Management Systems collect error log data directly from a printer output, receiving an alarm as the error occurs.

Fault management deals most commonly with events and traps as they occur on the network [Ref. 10]. Keep in mind though, that using data reporting mechanisms to actively report alarms or alerts is the best way to accomplish health checks of specific managed object's performance without having to greatly increase the amount of polling that must occur in order to ensure accurate error data is received.

(2) Configuration Management. Configuration management is probably the most important part of network management because a network cannot be accurately managed unless the configuration of the network can be managed [Ref. 10]. Changes, additions and deletions to the network configuration need to be coordinated with the Network Management Systems

(NMS). Dynamic updating of the configuration needs to be accomplished periodically to ensure the configuration is correct and known to the NMS.

(3) Accounting. The accounting function is usually left out of most implementations of NMS because LAN based systems are not known to promote accounting type functions unless dealing with hosts such as IBM mainframes or DEC VAX machines [Ref. 10]. Others rationalize that accounting is a server specific function and should be managed by the system administrator.

(4) Performance Management. Performance is a key concern to most MIS support people [Ref. 10]. Although, very important, it is considered difficult to be factual about some LAN performance parameters unless employing the Remote Monitoring (RMON) technology. A monitoring system is a cohesive, integrated set of elements for collecting information, offering the appropriate analysis and responding as needed. RMON agents can be a key element of the monitoring system.

RMON agents gather information for analysis tools. They are embedded in network products, such as hubs and switches, are available as stand-alone probes and in network interface cards. RMON agents also gather physical and data link layer data for a single LAN segment and collect data at the network and application layers for analysis of flows between parts of the enterprise network. Combining these agents gives detailed information about any LAN segment as well as end-to-end traffic analysis across complex networks. Although RMON agents are very useful, one should carefully weigh

what is pertinent to what can be accomplished in other ways without having to spend a bundle [Ref. 10].

Performance of Wide Area Network (WAN) links, telephone trunk utilization, etc., are areas that must be revisited on a continuing basis as these are easiest to optimize and realize savings.

Systems or applications performance is another area in which optimization can be accomplished. However, most network management applications do not address the issue in a functional manner.

(5) Security. Most network management applications only address security applicable to network hardware such as preventing an intruder from logging into a router or bridge [Ref. 11]. Some network management systems have alarm detection and reporting capabilities as part of physical security (contact closure, fire alarm interface, etc.) None really deals with system security as a function of system administration. This thesis will explore security policy management issues in later chapters.

(6) Chargeback. A chargeback system is intended to function as an IT Management control system by establishing the proper balance between controlling costs, stimulating use, and encouraging efficient use of IT resources. By making users who consume IT resources responsible for their costs, chargeback systems encourage more judicious use of IT resources, as well as promote a higher quality of service from providers [Ref. 9]. Chargeback has been done for years in the mainframe environments and will continue to be accomplished since it is a way to charge the end user for only the specific portion

of the service that he or she uses. Chargeback on Local Area Networks presents new challenges in that so many services are provided. In many implementations, chargeback is accomplished on the individual Server providing the service.

While chargeback is a challenge on broadcast based networks such as Ethernet; it is realizable on ATM networks that dynamically allocate bandwidth based on end users' needs. As technology associated with monitoring LAN and WAN networks evolves, chargeback will be integrated into more and more systems [Ref. 9].

(7) Systems Management. Systems Management is the management and administration of services provided on the network. A lot of implementations of NMS leave out this very crucial part. However, this is one of the areas in which Network Management systems can show significant capabilities, streamline business processes, and save the customer money with just a little work [Ref. 12]. There are many good Commercial Off The Shelf (COTS) products available to automate system administration functions and these products can be integrated into the overall NMS very easily.

(8) Cost Management. Cost management is an avenue in which the reliability, operability and maintainability of managed objects are addressed. This function is an enabler to upgrade equipment, delete unused services and tune the functionality of the Servers to the services provided. By continuously monitoring Mean Time Between Failures (MTBF) and Mean Time To Repair (MTTR) statistics, costs associated with maintaining the network as a system can be tuned [Ref. 16]. This area is an MFA that is driven by IT

management to ensure the highest degree of performance is received from the money allocated [Ref. 16].

c. Common Implementations

Most implementations of medium and large network management systems center around a Network Management Center (NMC) of some sort [Ref. 11]. From this location, all management related data is sent and processed. While several EMS are used to manage their specific areas, all of the data comes back to the Manager of Managers application, which runs on the NMC. Most fault detection, isolation and troubleshooting is accomplished in the Network Management Center and technicians dispatched when the problem has been analyzed and the solution devised. Several company locations may be involved in the overall network spanning thousands of miles around the globe (see Figure 3).

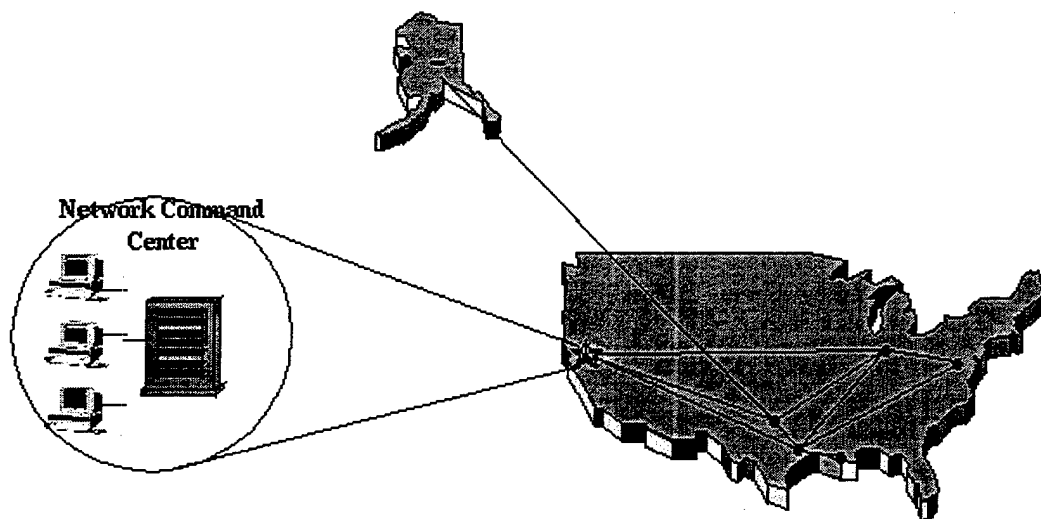


Figure 3. Network Management Center (each dot representing a MNC). [Ref. 8]

(1) Management Focus. The management focus for this scenario is on the Network Management Center driving the total operation. Detection, troubleshooting and dispatching are accomplished from the NMC. This operational focus is a carry over from the old Netview days when the center of the picture was a huge IBM Mainframe that did all the work. If a Network Management Center is not utilized today, consider what it will cost not only for the hardware and software, but the people to accomplish this and their level of expertise [Ref. 12].

(2) The Right Implementation. If MIS Managers, are looking at the benefits of network management to reduce downtime and overall cost, they need to make sure that the business case requirements drive the implementation and not the implementation drive the business cases [Ref. 12].

Systems integrators, need to make sure the requirements are collected before any implementation. When the requirements are put in place, it is the Systems integrators job to make sure management is informed as to what each implementation segment will cost and what that capability brings to the overall MIS function [Ref. 12].

d. Business Case Requirements

In today's world, any implementation must follow the business case associated with what will be implemented [Ref. 12]. The implementation must solve a business problem or increase efficiency of the current methods of

accomplishing work with lower costs. If the solution does not save time and money, it probably is not worth accomplishing.

(1) Definition. The hardest part of building a business case is the gathering of the information. One must define the problem at hand in a general sense and then look for specific problems that network management can address [Ref. 12].

The developer of the business case must look at the current way each section accomplishes its day to day work. The case for a NMS can be definitized by documenting current work processes that may be automated by the system.

The network manager should look for ways to save the organization money. Keep addressing the need to get the MIS into the organization to take advantage of the services they provide. A more efficient organization will be the result.

(2) Levels of Activity. There are four levels of network monitoring and management activity that one must understand before applying management to a specific service or device [Ref. 12]. These four levels of activity are as follows:

Inactive

This is the case when no monitoring is being done and if an alarm was received in this area, it would be ignored.

Reactive

The problem is reacted to after it has occurred yet no monitoring has been applied.

Interactive

Monitor the components but interactively troubleshoot to eliminate the side effect alarms and isolate to a root cause.

Proactive

This is where the system provides a root cause alarm for the problem at hand and automatic restoral processes are in place where possible to minimize downtime.

These four levels of activities outline exactly how a support organization is dealing with problems today and where an MIS manager would want them to be. Within the support organization are teams with different goals and focuses (i.e. Unix support, desktop support, network support, etc.). While a specific alarm may warrant an inactive approach by one team, it may demand a proactive approach from another.

(3) Today's Implementations. Of the network management implementations done today, very few really address the needs of the business. Most are implemented with good intentions but are focused away from increasing efficiency [Ref. 13].

In a multiple site network, technicians, engineers and support personnel are required at each major location. No one knows those local environments better than the people who work with them on a daily basis. No one knows the people of the organization better than the Help Desk staff because they are the first lines of communication between the people and the MIS support organization.

Network management elements are considered, among other things, tools used for troubleshooting accomplishment. The local support staff could benefit greatly from the use of this tool. As such, most implementations give read-only access to these systems. The ability to focus these tools at a local level is paramount to increasing the effectiveness to the local support staff. In some implementations, where read/write access is provided, it is accomplished through a (X-Windows) GUI, which does not work very well across low speed links [Ref. 12].

Most implementations focus these tools at a global level in that they are located in the Network Command Center. When a trouble ticket is generated from the NCC, it reflects a problem or symptom generated by the network management elements and/or the Manager of Managers. Sometimes, the local technician can not relate to this symptom because he or she does not

understand where this message came from or why. Without access to the management element and familiarity with the product, the technician usually starts problem isolation in a "cloud" looking for the problem [Ref. 13].

When a global problem occurs, in these scenarios, the information is concentrated and orchestrated by the Network Command Center. Additionally an outage can black out management of a geographic location when centralizing the management resources. Figure 4 illustrates how this occurs.

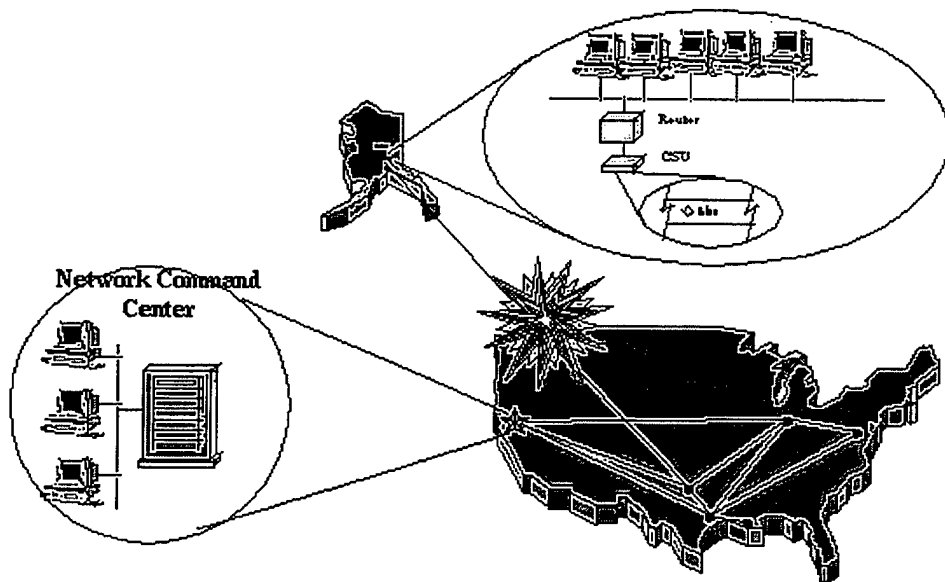


Figure 4. Network Command Center. [Ref. 8]

As far as the Network Management Center is concerned, all of the devices beyond the point of breakage are down. In fact, without alarm correlation, all of the devices will be depicted as bad. Even with alarm correlation, it can only be accomplished on one side of the link. No network management capabilities exist at the remote site to help troubleshoot the problem.

e. *System Focus*

The ideal network management system should be designed and implemented around the real work processes [Ref. 14]. It should focus the tools toward those staff members supporting the managed area in a manner which makes their job easier and faster. Information associated with a problem or symptom should mean something to the support personnel. If they see the problem at a glance, they should know which specific area the problem belongs to and what to do to get started isolating the problem. Other personnel in the organization should know that a technician is looking into the problem since the problem may affect other areas.

Help Desk personnel should know what is happening and who is working on what at a glance. If they are not familiar with the system in question, they should have adequate information at their fingertips to guide them in what to do, who to call, and what steps to take, even what questions to ask.

Additionally, lists of the problems that affect other sites should be available to those personnel at a glance. The information must be at the fingertips of the other sites' Help Desks so that they know, in near real time, what is going on.

In summary, the focus of information should be local when it is a local problem and global when it is a global problem. Figure 5 depicts a more distributed system providing global information with local focus. In this system, alarms can be passed from site to site and even around a problem with simple client-server database techniques.

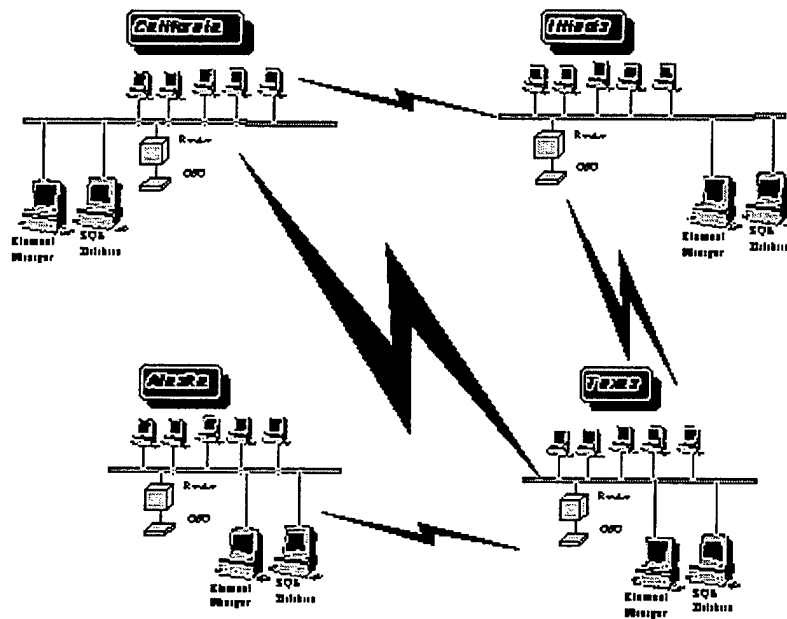


Figure 5. Distributed System. [Ref. 8]

In the scenario in Figure 5, if a link breaks, local tools and alarms are still available. Alarms concerning the overall health of other links and connectivity can be passed to other sites, even around a problem. Using a Serial Line Interface Protocol (SLIP) or Point-to-Point Protocol [Internet] (PPP), dial up link between management elements can be used to pass critical data about a link outage in near real time.

Network management across low speed wide area links does not really make sense [Ref. 8]. In this case, bandwidth is costly. Usually, there are the monthly charges for the links. Consider also that bridges or routers interconnect most WAN links. On the back side of these devices are networks capable of 10 Mbps, 16 Mbps or even 100 Mbps. On the link side 1.544 Mbps,

512kbps or even 19.2kbps links may be realized. Actual polling of network management elements (SNMP) could consume these links, drastically reducing the operational capabilities of the links. The question to ask is *do you want to increase the bandwidth across these links just for network management or do you want to distribute the management polling to local area concentrations and just pass the real alarm information?*

In summary, providing global information with local focus is the key [Ref. 8]. A more distributed system provides this while ensuring the availability of local tools, alarms and connectivity with other sites should a link fail.

3. Management Functional Domains (MFDs)

a. Introduction

Management Functional Domains are the segmentation of the Enterprise Network Management System into localized functional domains. The grouping of functions within specific domains allows alarm messages to be routed around problems especially when multiple usable paths exist. Furthermore, automated SLIP or PPP sessions will enable alarm passing through dial-up lines [Ref. 12].

Not just alarm messages need to be passed to other affected MFD's. Alarm correlation information and automatic diagnostics are examples of other information relative to a fault that provide a better picture of what is really happening on the other end.

By having the ability to validate the alarms around a broken link, one can quickly and efficiently determine the root cause. CPU utilization

associated with correlating the alarms is very low compared to the AI inference engine based alarm correlation. One simply looks for alarms that are common to both sides.

b. Building Requirements

The following is a list of steps to take to develop a requirement matrix associated with the management of network components and functions [Ref. 12].

- Develop a list of information attainable from each managed object. Describe in detail, each piece of information such as what the data element is, average versus actual, counter, raw integer or a text message.
- Take the list to the support organization responsible for that device function and have them decide what's pertinent to their way of doing business. Focus on information that will enhance their ability to accomplish their job.
- Formulate the reporting strategy for the device.
 - What elements of information are pertinent to alarm reporting?
Is there a real-time requirement for the reporting?
 - Establish thresholds, e.g. three counts in a one hour time period.
 - Establish the priority of the alarm and any thresholds associated with priority escalation of the alarm.
 - Establish any diagnostic processes that could be run automatically or the Help Desk could perform that would make their job easier.
 - Establish acceptable polling intervals (Every five minutes, ten minutes, one-hour, etc.)
 - What elements of information are pertinent to monthly reporting?
 - Availability of devices and services.

- Usage and load.
- What elements of information are pertinent to trending and performance tuning of network components and functions?
 - Look at ways to combine data elements or perform calculations on the data to make it more useful to the support organization.
- Interview Management to ensure the Network Management System is managing all areas pertinent to the business unit.
- Explain the role and objectives of the Network Management System.
 - Increase productivity throughout the support organizations.
 - Reduce the Mean Time to Repair times on the correction of problems.
 - Provide a proactive approach to the detection and isolation of problems.
 - Enable collaboration and the flow of information across support departments and sites.
- Gather the requirements for the management of any function important to the business unit.
 - Do not limit these functions to only SNMP manageable devices.
 - If the devices associated with a function have no intelligence whatsoever, go back to management later with a proposal to upgrade the devices.
- Go implement the requirements. Focus each implementation toward each requirement while integrating the total system.
- After implementation of each piece, notify the support organization associated with the managed object or system that monitoring has started.
- At the first reporting period, go back and revisit the requirements with each support organization and management.

- Reestablish requirements if necessary.
- The reports and types of data will change as each support organization becomes better informed.

During implementation, focus the alarm messages toward the Help Desk. They are the front line of any MIS organization. Keeping them well informed of problems is paramount to the successful deployment of the Network Management System [Ref. 12].

Perform "Dry Runs" of alarms and the diagnostic steps associated with getting the problem resolved in a quick and efficient manner. Have the appropriate support organizations participate so that all diagnostic steps can be identified and included. Train the Help Desk to input troubleshooting procedure pertinent to their function into the diagnostics table. This can include anything from a user calling in with an application problem (i.e. MS Word), to filling out forms for a specific service to be provided to an end user [Ref. 12].

The skills associated with the support organizations in one MFD may be different from another MFD. The gathering of diagnostic procedures allows a "sharing of the wealth" of knowledge across the enterprise. The diagnostics procedures are a knowledge base of information, by symptom, of problems and taking and what needs to be accomplished to correct the problem. Having the skills of Desktop Support, Unix System Support, Network Support, etc., at the fingertips of Help Desk personnel increases their ability to logically react to problems as they occur [Ref. 12]. Automatic diagnostics are also very important. This does not necessarily mean the procedure will take a device out

of service and run tests on the unit or circuit, although this may be acceptable in some instances. Polling or the gathering of information concerning the device in question and its configuration can be run without taking the device off-line.

The Network Management System, as a total integrated system, must be modular and easy to expand and contract as the needs of the business change [Ref. 9]. Element Management Systems, whether they are third party products such as SunNet Manager, HP Openview, Netview, NetMaster, 3M TOPAZ, Larsecom's Integra-T, or in-house developed pollers, need to be easy to integrate into the whole system [Ref. 13]. Recognize that in the architecture, no EMS is really aware of another. Awareness across EMS needs to be accomplished at a higher layer so that the EMS can focus on their area of management within their MFD.

Functions such as Alarm Correlation, Diagnostics across EMS, etc., can be accomplished using artificial intelligence principals within a relational database [Ref. 9]. Almost all Manager of Manager products employ an AI inference engine to calculate the probability that one component is so many percent more probable to break versus another. The inclusion of the AI inference engine drives up the cost because of the engine AND the hardware required to run these types of calculations. These types of decisions need to be accomplished through the support organizations within the MFD because they know the local environment better than any machine or personnel at another site [Ref. 9]. The overall application will serve its purpose better if it is more tightly integrated into the business units.

The application's of AI still needs to be applied but at a much different level. Network General Distributed Sniffer Servers are an excellent application of AI technology. By analyzing the relationships of protocols, traffic, connections and LAN control mechanisms, the Decision Support System uses AI to sort out problems at a very low level before these problems become user identifiable and cause degradation or downtime [Ref. 9].

Additionally, AI can be used to capture the heuristics of network behavior and help with the diagnostics. The knowledge available from past alarms and the problems associated with the isolation and correction of the alarm needs to be incorporated into the overall system.

4. A New Paradigm for Network Management

As increasing numbers of corporations deploy intranets for enterprise-wide connectivity, major network providers are strategizing new ways to enable MIS departments to leverage their intranets to manage enterprise networks [Ref. 20]. Called Web-based management (WBM), these strategies allow administrators to monitor and maintain their networks using the same World Wide Web functionalities that make intranets such effective communication tools [Ref. 20]. These functionalities permit administrators to use any Web browser on any network node to quickly and easily configure, control, and access networks and their individual components. WBM is a revolutionary network management solution that will transform how users manage their networks [Ref. 20].

This section examines the emergence of WBM and its major benefits, and explains the two main architectural approaches to WBM. It provides an overview

of emerging standards for WBM and a brief overview of WBM and the World Wide Web.

a. *The Advent of Web-Based Management*

WBM is a result of the growing popularity of intranets [Ref. 21]. Intranets are, in effect, private World Wide Webs. They are increasingly used as a primary means for sharing information within an organization. Intranet networks run Transmission Control Protocol (TCP) and are isolated from the external Internet by security firewalls. They are constructed with Web-enabled servers using protocols related to the Hypertext Markup Language (HTML). Intranet users communicate with servers using friendly, easy-to-use Web browsers from any networked platform or location. Connectivity is simple, inexpensive, and seamless. Secure remote intranet access is typically achieved today using direct dial-up capabilities, although new techniques such as Point-to-Point Tunneling Protocol (PPTP), pioneered by a handful of vendors, are providing new secure intranet access options via the Internet [Ref. 20].

With intranets in place, administrators are finding other benefits of Web technologies. Many, for example are rethinking traditional client/server models in order to further optimize network usage and reduce development, equipment, and support costs [Ref. 21]. Since Web browsers require only moderately powered machines with modest disk space, administrators can shift much of the computing and storage tasks to Web-based servers and allow clients to rely on simple, inexpensive computing platforms to access them. This "thin client/fat server" model lowers hardware costs and offers users greater flexibility.

Web technologies are accelerating rapidly, in no small part because of the competitive nature of this industry. The battles rage on many fronts. Netscape and Microsoft, for example, are rapidly introducing new Web browser and server capabilities in their struggle to dominate the browser, server, and groupware software markets [Ref. 20]. A host of vendors, including Sun Microsystems and Oracle, are aggressively promoting the thin client/fat server model as a means to dislodge Microsoft and Intel from their dominant positions. Another battle rages to set standards in electronic commerce for Web-traffic encryption [Ref. 20].

Still another clash for market leadership centers on emerging "channel" or "Web broadcasting" technologies, with combatants such as Marimba and BackWeb competing fiercely. In the network management arena, major players—including IBM/Tivoli, Sun, and Hewlett-Packard—are racing to offer management platforms enabled with Web technologies [Ref. 21].

b. The Benefits of Web-Based Management

WBM merges Web functionalities with network management to provide administrators with capabilities beyond traditional tools. Since administrators using WBM can monitor and control enterprise networks with any Web browser at any node, they are no longer tied down to management workstations and can eliminate many interoperability issues that arise with multiplatform structures [Ref. 21]. WBM provides graphical interfaces that present information in a more visual and useful fashion than conventional, command-driven TELNET screens. Browser operation and Web page interfaces

are known paradigms for today's users of the World Wide Web. As a result, both reduce the costs of training MIS personnel and enable a wider range of users to utilize network status information.

In addition, WBM is an ideal means for distributing information about network operation [Ref. 20]. For instance, by directing their browsers to a designated intranet Web site, users can access network and service updates, thereby eliminating calls to enterprise help desks or MIS staff. Moreover, since WBM requires only the installation of a Web-based server, integrating it into intranets is a quick and painless task.

c. Two Strategies for Implementation

There are two basic approaches to WBM [Ref. 21]. They are evolving in parallel and are not mutually exclusive. The first is a proxy solution in which a Web-based server is added to an intermediate station (the proxy), which in turn communicates with end devices (Figure 6). A user at a browser communicates with the proxy using the Web's Hypertext Transmission Protocol (HTTP), while the proxy communicates with end devices using SNMP. Typically, vendors develop proxy solutions by adding a Web server to an existing management product, such as 3Com's Transcend® Enterprise Manager. This lets them leverage product capabilities such as database access and SNMP polling.

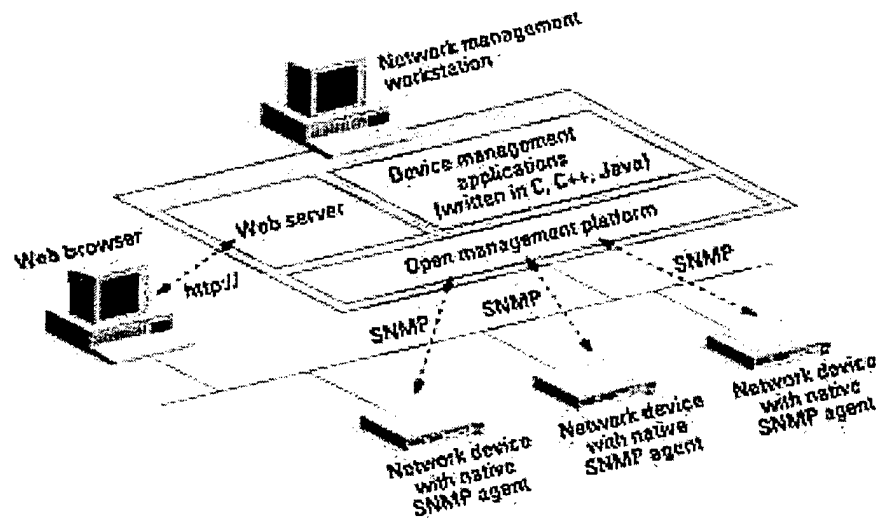


Figure 6. The Proxy Solution for WBM. [Ref. 20]

In the second WBM mode, the embedded approach, a Web server is actually embedded in the end device. Each device has its own Web address, and an administrator simply visits a device's address with a browser to manage that device (Figure 7).

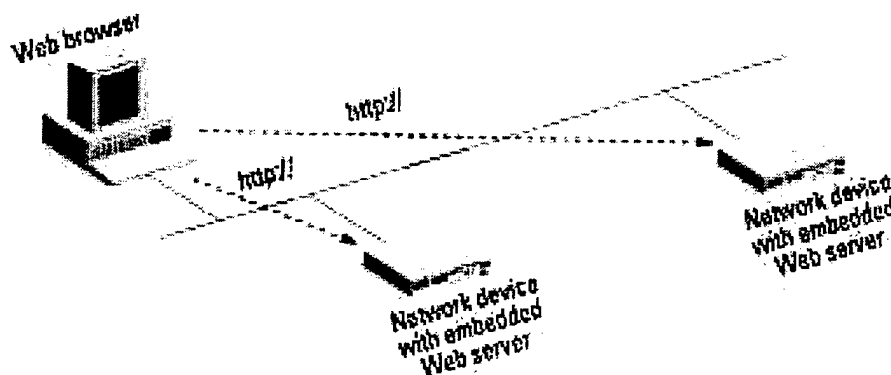


Figure 7. The Embedded Approach to WBM. [Ref. 20]

The proxy approach preserves all the advantages of today's traditional workstation-based management systems and products, with the added benefit of flexible access. Since the proxy communicates with all network devices, it can provide an enterprise view of a network's physical devices and

stations as well as its logical entities like virtual LANs. Also, because the proxy-to-device communication remains SNMP, this solution works with SNMP-only legacy devices. It does, however, require a workstation to operate.

The embedded server approach, on the other hand, brings graphical management to individual devices. This promises to be a much simpler, easier-to-use interface than today's command line or menu-based TELNET interfaces [Ref. 21]. The Web interface will deliver easier operation without sacrificing functionality.

In the enterprise networks of tomorrow, it is likely that both proxy-based and embedded Web server capabilities will be used for management [Ref. 21]. Large organizations will continue to need enterprise-wide monitoring and management available only with a proxy solution. Also, a proxy solution can manage the large installed base of native SNMP-only devices. At the same time, embedded Web servers will benefit enterprise networks because of their greatly improved interface for setting up and managing new devices.

d. Emerging Standards

Open standards are essential to reducing the complexity and costs of network management, and two WBM standards are presently under consideration [Ref. 22]. One is the Web-Based Enterprise Management (WBEM) standard effort launched in July 1996. WBEM is an initiative driven by Microsoft with support from more than 60 vendors. The proposed standard is an object-oriented means to abstract management data which today is gathered from multiple sources (such as devices, systems, applications) using multiple

protocols (for example, SNMP, Desktop Management Interface (DMI). WBEM can consolidate the information and management capabilities and expose them through a single protocol. WBEM is positioned as "embracing and extending" current standards such as SNMP, DMI, and Common Management Information Protocol (CMIP), not as a replacement.

WBEM places more emphasis on enterprise management than on the fact that it is Web based. Although WBEM lends itself to a Web implementation, its real goal is consolidated management of all network elements and systems. This includes networking equipment, servers, desktops, and applications. Key to WBEM is a new protocol, Hypermedia Management Protocol (HMMP). This transport protocol handles such functions as retries, packet pacing, and confirmation of delivery, and allows for a message to consist of a stream of several packets [Ref. 22].

The other proposed WBM standard is Java-Management Application Programming Interface (JMAPI), championed by Sun as part of its Java Standard Extension Active Program Interface (JSEAPI) framework. More than its name implies, JMAPI is a complete network management application development environment that provides a comprehensive list of features that developers had previously been forced to use together [Ref. 22]. These include user interface classes for creating property sheets, tables, and graphs; network APIs for SNMP; remote procedure call frameworks; database access methods; and style guides. In theory, JMAPI-compliant applications will interoperate smoothly using the same look and capabilities all over the Web.

At present, these particular standards are far from adoption. WBEM proposals today require significant additional writing and prototype "proof of concept" work before moving forward. WBEM is looking to the Internet Engineering Task Force for open standards support, but so far the proposals and refinement process remain very Microsoft NT-centric [Ref. 22].

Networking vendors have expressed concerns regarding the impact WBEM and HMMP could have on network performance and the costs of implementing them for networking devices. JMAPI is less ambitious but also has holes, such as a lack of specific documentation on how it is to interface with SNMP. JMAPI so far has been very Sun-specific without being driven as a true open standard. These concerns must be addressed before one or both are deployed for networking devices and WBM solutions.

e. *Web-Based Management and the World Wide Web*

WBM solutions will fall short of their potential if they do not transparently leverage the Internet and the World Wide Web for information and product delivery [Ref. 22]. While vendor Web capabilities are generally considered independent of WBM, the two are complementary and should be tightly coupled. What better place for a network administrator to access information on a vendor's Web site than from the vendor's browser-accessed WBM application?

Network administrators using WBM should have "fingertip" access to the information they need to run their networks. Such information could include complete product documentation, release notes, operation tips,

frequently asked questions, technology tutorials, training materials and schedules, new product data sheets, and order status. The Web, together with WBM, is an ideal vehicle for delivering software-based products, including new application releases under subscription, patch releases, and new agent releases. When combined with WBM capabilities to survey a network for out-of-release products, such software distribution can be semi-automatic, greatly simplifying the administrator's task.

5. Questions to Ask

As a MIS Manager, when staff or vendors concerning Network Management approach, there are a few key questions to ask [Ref. 17].

a. How Much will the System Cost?

A salesman specifying the system to the MIS Manager dictates a lot of systems implemented today. Salespersons typically push huge amounts of hardware and software towards the problems at hand. Some vendors will even say that cost is not important; it's the capability that counts. Additionally, because a network management system must be customized to the local environment, there are a lot of hidden costs beyond the hardware and software.

b. Will the Proposed System Integrate into and Enhance my Current MIS Support Capabilities?

A lot of MIS Managers may be shortsighted in that they are not demanding that the overall system be tightly integrated into the business units

[Ref. 18]. If the system serves no business purpose, technology is being bought for technology's sake... the system may be doomed to failure.

c. *Is the proposed system modular in design?*

If everything in a Network Management System is loaded on one box, there is the potential for inefficient use of computing resources. If the system contracts, the one box may be under-utilized; if it expands, a bigger one may be needed... potentially losing money every time [Ref. 18].

d. *Is the product proposed just an Element Management System or is it an Integrator of Element Management Systems?*

MIS Managers may be sold a product like HP Openview or IBM Netview 6000 as a Manager of Managers System. Although, some integration functions are capable in these systems, their ability to perform real work... like polling and gathering information is degraded [Ref. 18].

e. *What does the system monitor?*

Match the capabilities of the proposed Network Management System to the key I/T services provided. If it is not a good match now, it won't be later.

f. *Does the proposed system enhance the capabilities of the current support staff or does it add more support staff?*

Some systems will do nothing to enhance current support staff capabilities and add five or ten more personnel to the staff and add to the budget. Not to mention, these people are usually highly skilled specialists in Network Management... which don't come cheap.

The total picture of the entire enterprise should be looked at. Match what is proposed to what's currently operational. Ask the same questions for each site.

6. Conclusion

There are a lot of excellent products available today that provide capabilities to manage not just hardware, but services and applications. The ways these systems are implemented are also critical. Each management capability installed must match a business need for such a system. Additionally, these diverse systems must be integrated together and into the support organizations to achieve maximum effectiveness [Ref. 18].

B. POLICY-BASED MANAGEMENT

This section will explore the basics of Policy-Based Management in function and form. The reader should gain an appreciation for where this management technique fits in to his or her network.

Management involves monitoring the activity of a system, making management decisions and performing control actions to modify the behavior of the system. Management policy guides the decision making process. Policy is information, which influences the behavior of the components in the managed system [Ref. 17]. The size and complexity of large networks and distributed systems have resulted in a trend towards automating many aspects of management into distributed components. If policies are coded into manager components they become inflexible and cannot be reused in different environments. There is thus a need to be able to specify, represent and

manipulate management policies without building them into managers or manager agents. This permits manager objects to be reused in different environments and to be provided with the specific policies for each and enables dynamic changing of management policies for a system without changing the managers. [Ref. 17]

1. Definition of a Network Policy

A *policy* is a statement of a definite course of action, and defines the roles, obligations, and rights of constituents; for example, Policy 1 states a course of action:

Policy 1. Medical Records Shall Be Maintained Indefinitely

This policy asserts an obligation but does not assign this obligation to a particular role; nor does it assign a specific right for carrying out the obligation. Policies are intended to influence the actions of an enterprise's constituents [Ref. 17].

Different types of policies exist. For example, *Black's Law Dictionary* [Ref. 17] defines *public policy* as follows:

Public policy. That principle of law, which holds that no subject, can lawfully do that which has a tendency to be injurious to the public or against the public good. The principles under which the freedom of contract or private dealings is restricted by law for the good of the community. The term "policy," as applied to a statute, regulation, rule of law, course of action, or the like, refers to its probable effect, tendency, or object, considered with reference to the social or political well-being of the state.

In addition to classifying policies according to the type of enterprise that promulgated them or the type of constituents they apply to, the subject matter of

policies can also serve as a basis for differentiating between policies. For instance, policies that treat the protection of information from unauthorized access are generally termed *security policies*. According to Pfleeger [Ref. 17],

A [security] policy statement should specify the organization's goals regarding security (for example, protect data from leakage to outsiders, protect against loss of data due to physical disaster, protect the integrity of data); where the responsibility for security lies (for example, with a small computer security group, with each employee, with relevant managers); and the organization's commitment to security (for example, dollar expenditures, personnel assigned to task).

This definition is congruent with that given for security policy in the *Dictionary of Computing* [Ref. 17]. Policies specify the roles, obligations, and rights of both animate and inanimate objects. Policies form a natural hierarchy; for instance, the preceding policy could be decomposed into the following policy:

Policy 2. System administrators are responsible for assuring the resiliency of database transactions against system and media failures.

In contrast to Policy 1, Policy 2 *assigns* the obligation to a specific role; in other words, when a constituent assumes the role of system administrator, the constituent incurs an obligation to assure the "... resiliency of database transactions against system and media failures." Policy 2 could be further decomposed into Policy 3.

Policy 3. Database administrators of medical record databases shall implement logging and checkpointing procedures to protect against the loss of database transactions due to system or media failures.

Policy 3 differs from Policy 2 in that the policy scope is explicitly stated; that is, the policy applies only to a subclass of the class of database

administrators: those that administer medical record databases. *Policy scope* refers to the range of roles, obligations, and rights that a policy covers. *Policy domain*, on the other hand, defines the animate and inanimate objects in the environment that are constituents of an enterprise. For instance, if Policy 3 is that of a hospital, then animate (e.g., members of the hospital's data processing staff) and inanimate (e.g., computer processes executing logging and checkpointing instructions) objects that perform the role of database administrator at the hospital are *ipso facto* constituents of the hospital.

Moreover, policies are embedded in information systems as procedures. *Procedures* implement policies and take the form of artifacts. An *information system artifact* is a representation of a policy as a procedure, for example, a requirement, design, or system prototype. Thus Policy 4 could be translated into the procedure shown in Figure 8.

Policy 4. A "continuous" archive of all logging and check-pointing operations shall be maintained.

Procedures consist of "... rules and the ways to make exceptions to [rules] ...". Policies are often articulated as natural language statements. These tend to be vague; for instance, the meaning of the term "logging" in Policy 4 can be overloaded (i.e., have more than one meaning). The procedure in Figure 8, in contrast to Policy 4, is a more precise statement of what is meant by logging and check-pointing operations, and implements the policy as a set of rules. The procedure could in turn be decomposed into lower level procedures; for example, translated into machine-executable instructions embedded in an automated information system. Procedures thus form a hierarchy.

```
if a block, B, of the log, L, is written to secondary memory,  
then write B to the archive, A.  
else  
if the value of the item in B has been written at least once in L since the last  
checkpoint,  
then write B to A.  
else  
    if B was in main storage at the time the current checkpoint operation  
    started,  
    then write B to A.
```

**Figure 8. A Procedure for Continuous Archival of Logging and Check-
Pointing Data. [Ref. 17]**

Similarly, the volition of an animate object to perform an action can be embedded in an inanimate object as a procedure; for example, a computer-based multi-level secure database management system is "unwilling" to permit a subject (i.e., user or the computer process acting as a proxy for the user) to access classified information in a database unless the subject has been granted proper access rights to this data and has entered the correct password. This is a security policy of an enterprise represented in a computer program. Unfortunately, a programmer can create policy while coding a system. For instance, a programmer can implement logic that locks a user out of a system if

the user has failed to correctly enter a recognized password after four attempts; however there may be a corporate policy that limits this to three times.

A policy differs from a *requirement* in that the latter is a high-level procedure, which specifies how policies are to be implemented as procedures. For example, the *Trusted Computer System Evaluation Criteria* (TCSEC) [Ref. 17] consists of *security requirements* for implementing *trusted information systems*. The TCSEC requires among other things that trusted systems, certified at the B1 (or higher) level, enforce mandatory access control (MAC) policies; that is, there is a security requirement that MAC policies be embedded as procedures in a B1 (or higher) trusted system. The security requirements in the TCSEC were derived from national security policies, which are a specialization of public policy.

In summary, policies are composed and embedded in information systems as procedures, conflicts can exist between composed policies, and policy inconsistencies can become embedded in procedures.

Policy-based network management in the form of prioritization can control which network users have access to network bandwidth according to predetermined corporate policies and the amount of bandwidth provided to the user or workgroup based on their location, the application(s) they are using, their position, or other parameters; thereby controlling network performance in addition to network access [Ref. 17]. This type of mechanism can prevent or reduce congestion on a network running at near capacity levels. Today, enterprise networks provide the foundation for business. The advent of

networked business applications such as intranets, extranets, and electronic commerce over the Internet has created a boom in productivity and fostered unprecedented global commerce opportunities.

However, as more business applications are deployed over the network, which is clearly the industry-wide trend, the business becomes dependent on an intelligent network. As a result, enterprise network managers require the capability to control the use of the network and to allocate and prioritize network resources for different applications and user groups. This thesis will explore and categorize several network policies.

2. The Need for Policy Networking

Network managers need policy control over bandwidth-hungry applications, which consume bandwidth at the expense of performance and drive up the cost of expensive wide-area resources [Ref. 16]. Misbehaved applications can potentially shutdown the business. Additionally, there is an increasing need to manage security policies, the mis-managament of which can have devastating consequences.

The need for policy networking (see Figure 9), is being driven top-down by business requirements dictated by market and competitive forces [Ref. 16]. The network manager needs to be able to map these business requirements into specific policies that link the business needs with the desired network behavior. For example, if an organization is running an Enterprise Resource Planning (ERP) application (e.g. SAP R/3, Oracle Financials, Baan, Peoplesoft), for strategic competitive advantage, a policy can be established that gives ERP

traffic priority to network resources. The business policy is automatically translated into network behavior, such as Quality of Service (QoS) mechanisms, to prioritize ERP traffic ahead of other traffic.

The intelligent network provides a rich set of QoS and security mechanisms to enable business applications. However, utilizing these features can be a complex exercise for network managers. There is a real need to provide dynamic and automatic configuration of features in the intelligent network. Network devices must be dynamically tuned to support increased user mobility and new classes of applications such as Internet webcasting and multimedia applications that support data, voice, and video simultaneously [Ref. 16].

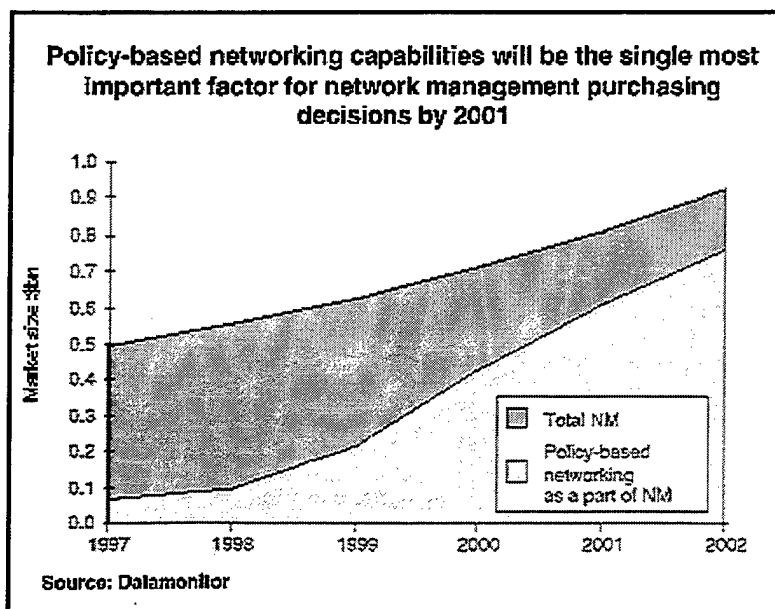


Figure 9. Policy-Based Network Management Growth Policy Networking Building Blocks. [Ref. 12]

Policy Networking enables business users and applications to use the intelligence that is embedded in a network. Simply put, Policy Networking makes

it easier for a network manager to take advantage of distributed network intelligence features.

The Policy Networking architecture is based upon four building blocks:

- **Intelligent Network**

Intelligent network devices---that is, routers, switches, and access servers--enable and enforce policy services in the network.

- **Policy Services**

Policy services translate business requirements into network configurations and activate policies for quality-of-service (QoS), security, and other network services.

- **Registration and Directory Services**

Registration and Directory services provide the dynamic binding between addresses, application profiles, user names, and other information data stores.

- **Policy Administration**

Policy administration provides the capability to centrally configure rule-based policies that control services within the network infrastructure.

3. Policy Networking Tree of Variables

POLICY-BASED NETWORK MANAGEMENT

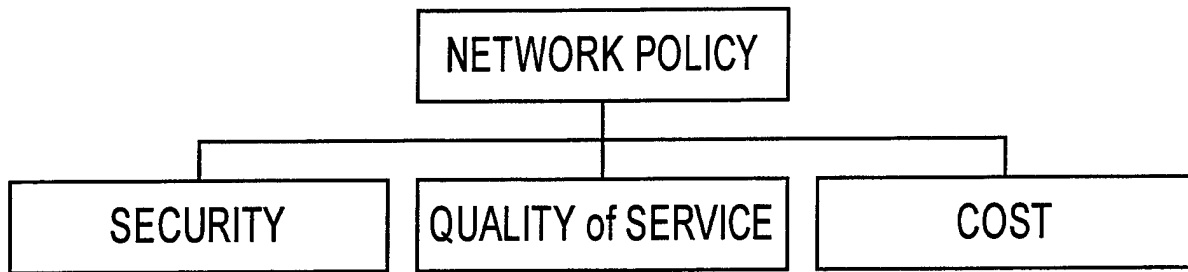


Figure 10. Network Policy Tree.

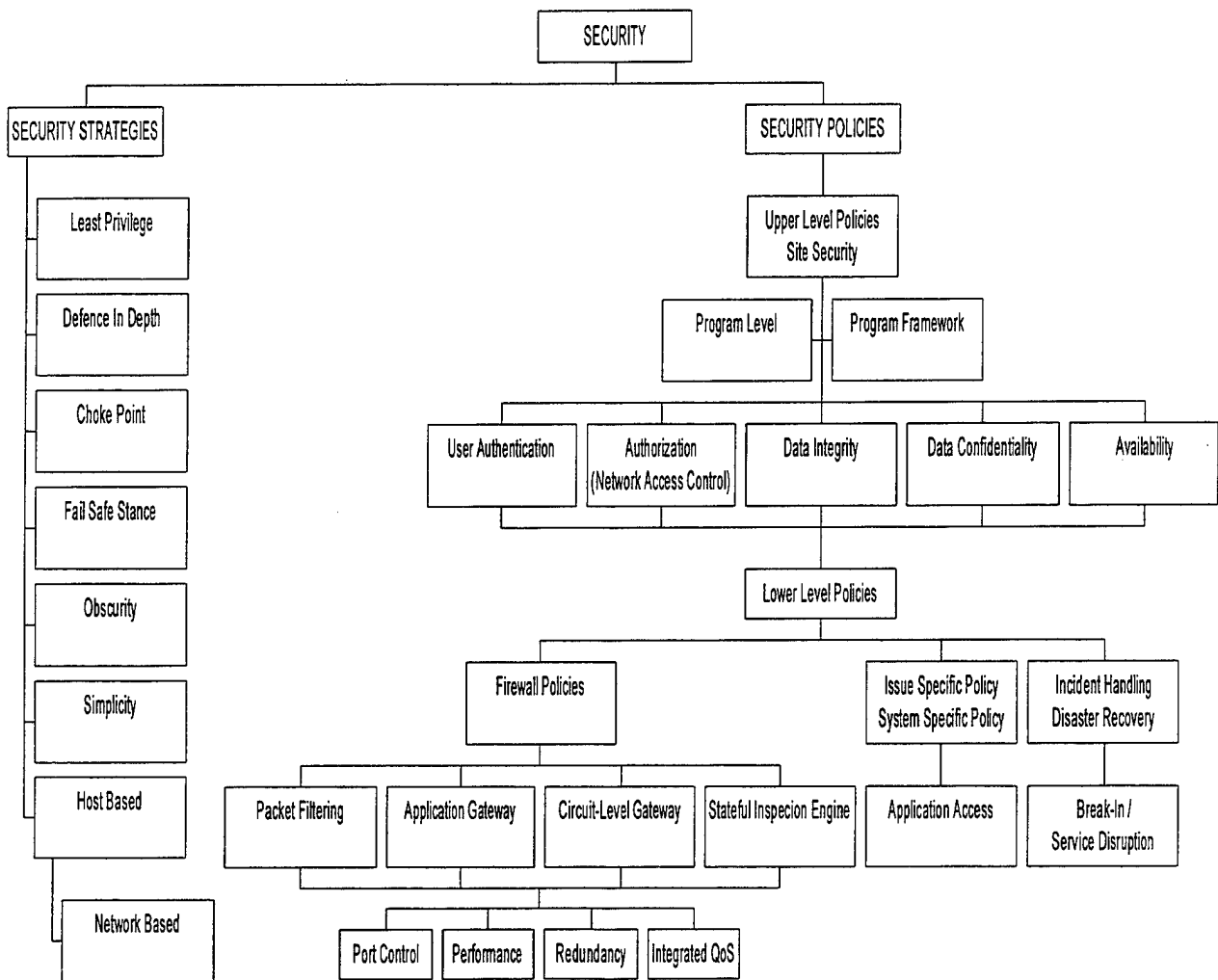


Figure 11. Security Policy Tree.

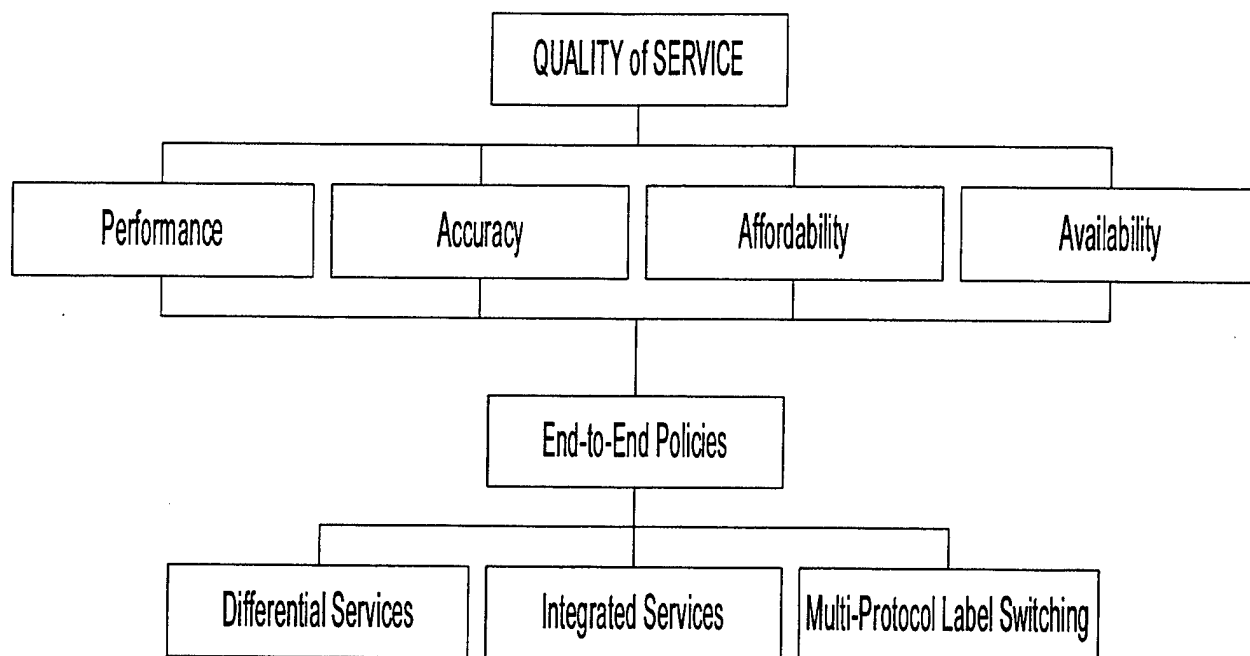


Figure 12. QoS Policy Tree.

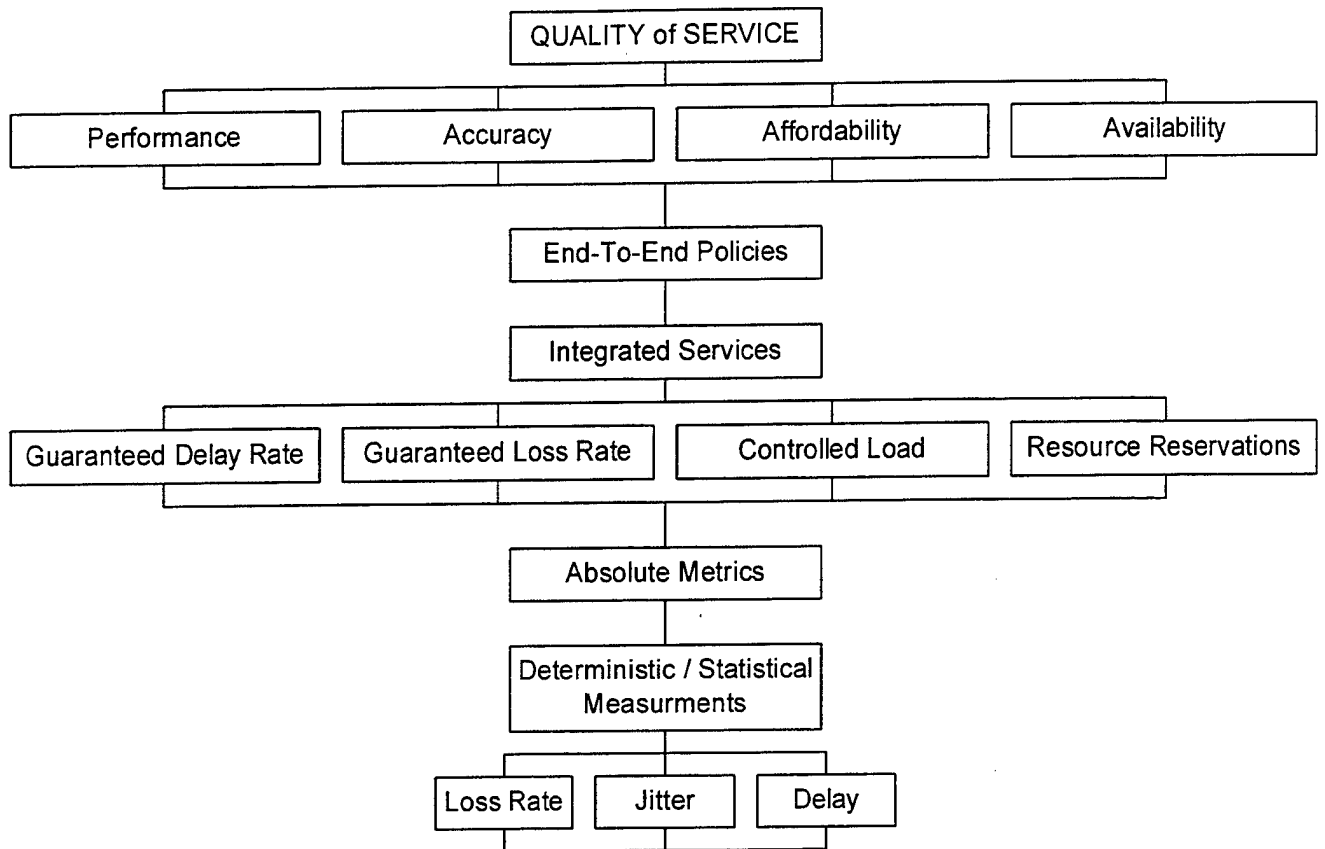


Figure 13. QoS Integrated Services.

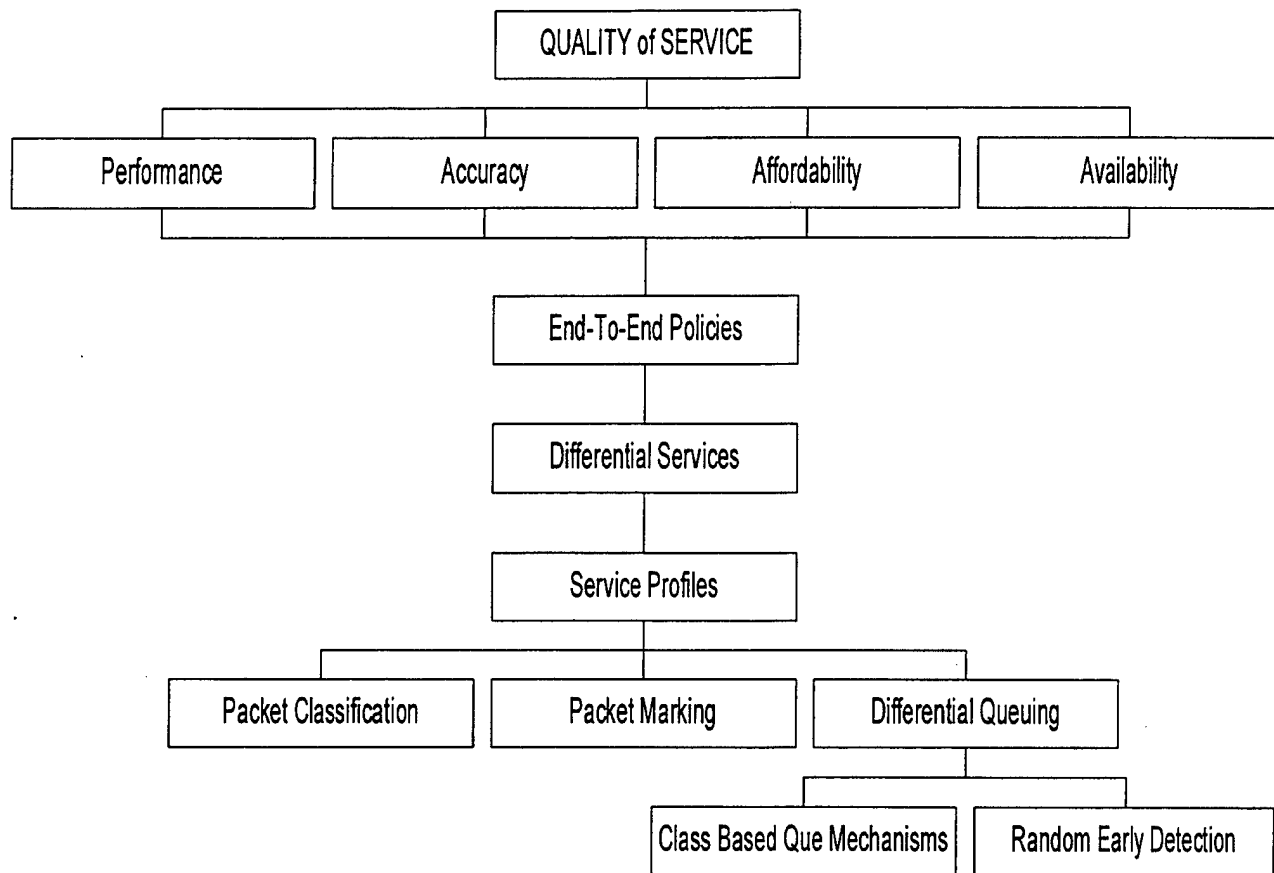


Figure 14. QoS Differential Services.

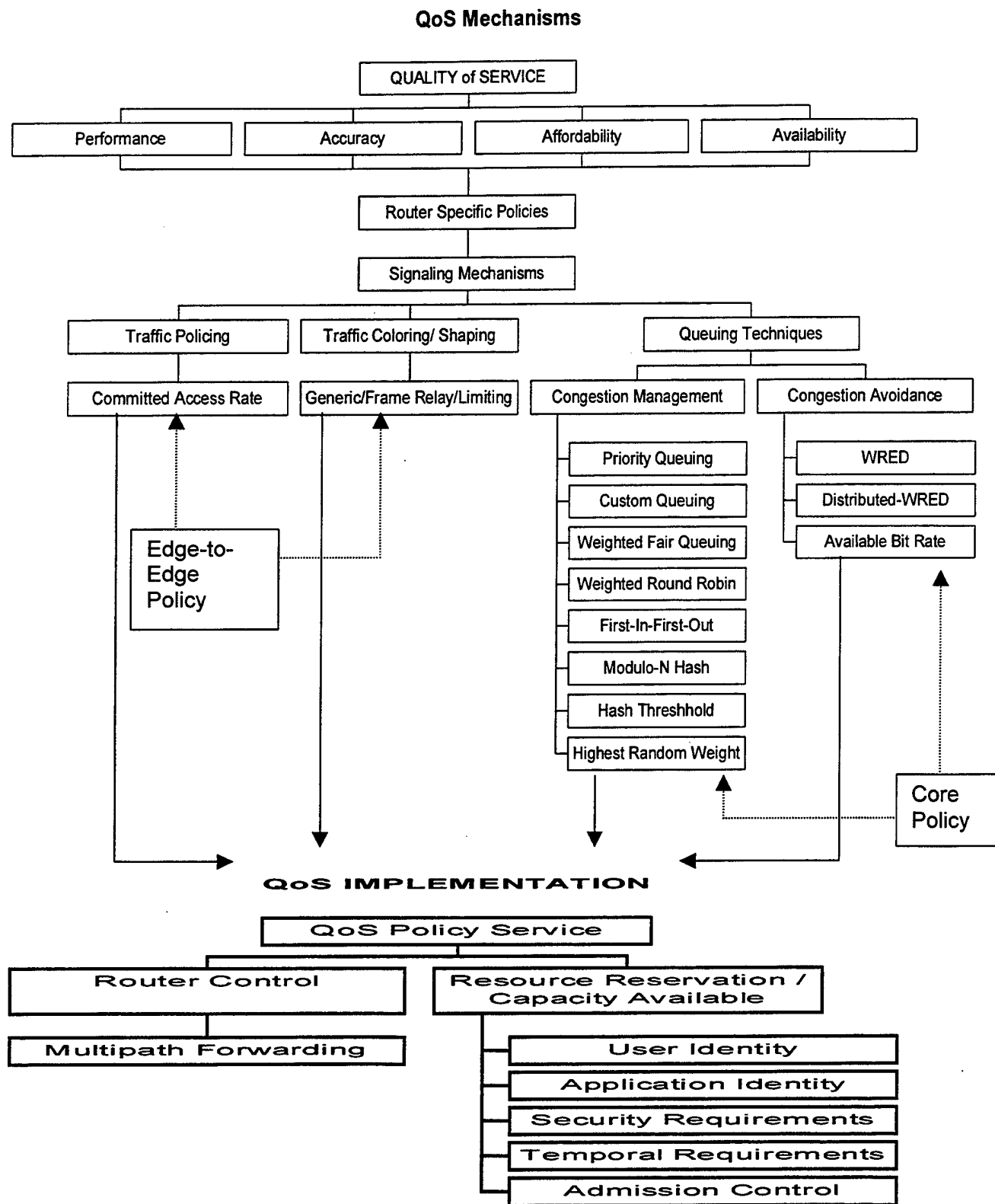


Figure 15. QoS Policy Mechanisms.

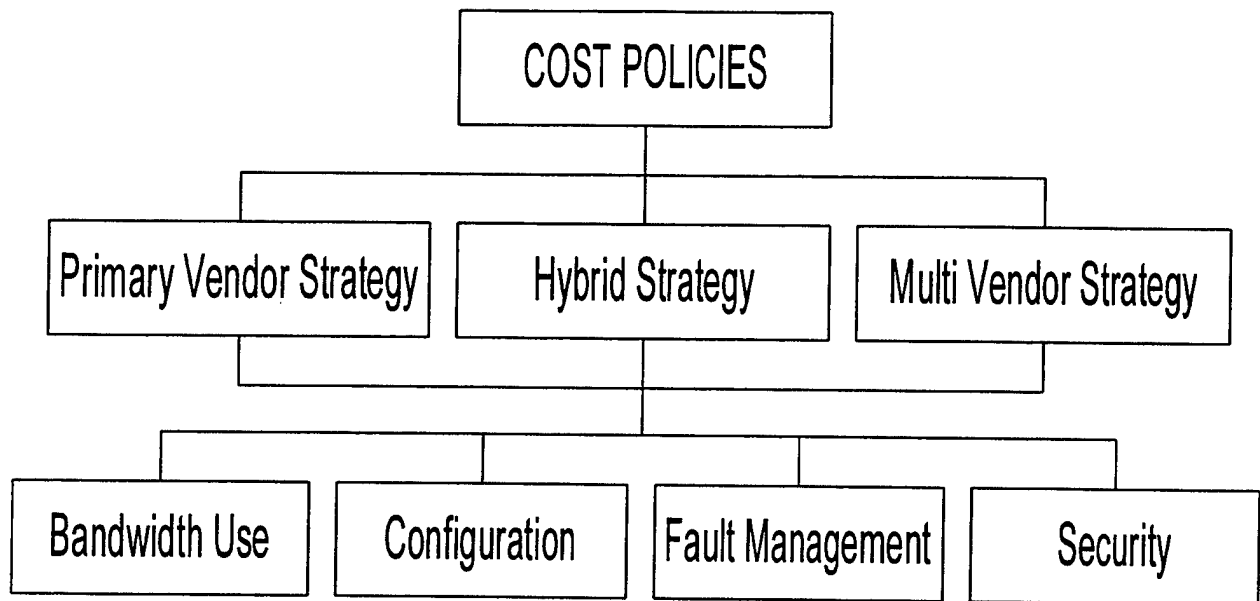


Figure 16. Cost Policy Tree.

THIS PAGE INTENTIONALLY LEFT BLANK

III. MANAGING QUALITY OF SERVICE

A. INTRODUCTION

Now that a background in management, network management and network policy basics has been discussed, the relatively new topic of network Quality of Service (QoS) will be explored. It should be a goal of the network manager to select a network policy or policies that will guarantee high QoS [Ref. 18]. This chapter provides an overview of a QoS architectural framework that explains how QoS applies in the network, and it provides details on technologies that typical QoS software provides in each piece of the architecture. The chapter concludes with some examples of how these pieces work together to provide QoS services that help realize the most from scarce network resources. Definitions of some QoS common terms are provided in the glossary.

QoS is an important aspect of the over all network management scheme. Networking users in general span three major market segments: major enterprises, network service providers, and the small and medium-sized business segment [Ref. 19]. Each segment has its own QoS requirements, but they also have many overlapping needs.

Network managers in today's enterprise networks must contend with numerous and diverse system platforms, network architectures, and protocols [Ref. 19]. Providing end-to-end QoS solutions across the various platforms often requires more than just linking them together; it also requires a different approach for each technology. Enterprises are increasingly depending on their

networks to carry complex mission-critical applications and databases such as SAP, PeopleSoft, and Oracle. These networks are also experiencing increased traffic from Web and multimedia applications [Ref. 19]. QoS prioritizes this traffic to ensure that mission-critical applications get the service they require, while simultaneously servicing these newer multimedia applications.

Internet Service Providers (ISPs) require assured scalability and performance [Ref. 18]. The ISP marketplace is also highly competitive and characterized by phenomenal growth. ISPs, who have traditionally offered best-effort IP connectivity, are now planning networks to transport voice, video, and other real-time, critical application data. ISPs need a model that will allow them to offer differentiated services to their customers, yet allow them to remain profitable. QoS provides the basis for a new business model by allowing ISPs to differentiate traffic from various customers or applications [Ref. 19].

In the small and medium-sized business segment, managers are experiencing first hand the rapid growth of business on the Internet. Not so long ago, the "global networked business" was just a concept. Every day we witness more and more businesses participating in the reality. Besides the increased demands of Internet traffic, small and medium-sized business networks must also handle increasingly complex business applications. QoS policies let the network handle the difficult task of utilizing an expensive wide-area network connection in the most efficient way for business applications.

Over the last 10-20 years, as global competition has increased, managers have increasingly sought ways to improve their organizations' competitiveness

[Ref. 19]. Often, as a new idea appeared, it was adopted rapidly, and some would say unquestioningly. All too often, these new ideas have ended up merely being short-lived fads that have produced few tangible benefits. Many of these ideas have been held up to ridicule in the popular cartoon, "Dilbert."

Managers of Information Technology organizations (IT) have proven to be susceptible to this same phenomenon. The frequent misuse of popular technology has been rampant [Ref. 19]. Currently, as suggested in the Introduction to this thesis, there appears to be the beginning of a groundswell of interest among IT managers in quality of service management (or service level management) [Ref. 19]. The question must be asked. "Is quality of service management another management fad, or is it a legitimate business focus?" Many cynics would be quick to say it is a fad. However, a review of today's situation quickly refutes such skepticism.

Traditionally, IT managers have measured the effectiveness of their organization by looking at the various hardware and software components for which they are responsible [Ref. 19]. The units of measure that were applied have been varied. The most common have been such things as availability, utilization, performance, etc. This has given IT a fragmented, misleading perspective of its effectiveness. It is a perspective that is distorted and misleading [Ref. 19]. It has led IT managers to believe that they have been doing an excellent job at meeting the needs of their clients for IT services.

In recent years, while IT managers have been congratulating each other on the ever-improving quality of service that they have been providing, their

clients have had a radically different perspective [Ref. 21]. They have been complaining about the inadequacy of the service and the unresponsiveness of the IT organization to client requirements. In recent years clients of the IT organization have become more and more dependent upon the services provided by IT [Ref. 19]. The very existence of many businesses depends upon the delivery of those services. In addition, their clients (users) have become more technically sophisticated. This has given them a better understanding of what they can reasonably expect IT to deliver. Together, these factors have led IT's clients to demand a much higher quality of service from IT. In some cases, it has also led the users to stage a coup, toppling the IT organization from its pedestal and replacing it with an out-sourcer's organization. The clients of IT hold the corporate power and control the corporate purse. One way or another, they are determined to receive the level of IT services that they require [Ref. 21].

Today's interest in managing the IT organization in terms of the quality of service it delivers has come to the forefront due to the convergence of two major factors: Supply and Demand. The confidence of IT managers has been shaken by the complaints of disgruntled clients and also by continuing reports of outsourcing deals [Ref. 19]. IT, and their clients, has a need to accurately measure and report on the quality of the service being delivered by IT. This is the demand side of the equation. True end-to-end service level measurement has only recently become practical. There are now tools emerging to support the measurement of the quality of the service being delivered by IT.

It must be remembered that not only is IT a service provider, it is also a major consumer of services. If those providing services to IT do not deliver a consistent level of service, IT will not be able to be consistent in the quality of services that it delivers to its clients. Those providing services to IT include groups internal to the company (e.g. data base administrators, application developers, etc.) as well as some external groups (e.g., telephone company). However, it must be realized that managing the quality of service delivered is not limited just to IT organizations [Ref. 18]. Certainly, for most organizations the IT function is a critical service provider, but there are many others. For example, the telephone company that provides voice and data links to other companies, is a major service provider. It is only reasonable to demand a guarantee of the quality of service being provided.

Is the current desire to establish effective measurement and reporting of the quality service being delivered a fad? Absolutely not! As will be seen in this chapter, quality of service management (service level management) is absolutely essential as a survival strategy for the IT organization.

B. DIVERGENT VIEWS

What has gone wrong? Why have the views of IT and their clients about the quality of service being provided by IT become so divergent? The bottom line answer is that the two groups do not have a common understanding of what is an adequate level of service or even about what the key indicators are of the level of service being provided [Ref. 19].

Until now, most IT organizations have not been measuring the quality of the service being delivered to their clients. Instead of measuring the quality of the service being delivered, IT has been measuring various attributes of the hardware and software used to deliver the service to its clients [Ref. 16]. Certainly this data is valuable for evaluating the performance of the IT employees or planning for adequate equipment capacity in the future. Technical metrics are also essential in assessing the quality of the service being delivered by the external network service providers (i.e. telephone companies). However, this technical data is not indicative of the level of service being delivered by IT to its clients.

There have been several flaws in the traditional strategy employed by IT [Ref. 19]. First, it has not been possible to capture equivalent information for every component. That is, in examining the end-to-end quality of service being delivered, there are numerous components that impact that service. There are various sources of information about those components and about some portion of the service being delivered. These different data sources do not capture the same types of information, nor is their data capture synchronized. The result is that the data available has been fragmented, scattered and disjointed.

Second, for some components, only limited data has been available [Ref. 19]. In some cases, no data relevant to quality of service has been available. The third flaw has been the type of information captured. Certainly, Central Processing Unit (CPU) uptime, dropped packets, network congestion, etc. will each impact the quality of the service that is delivered. However, these factors

can only indirectly suggest something about the quality of service being delivered. Finally, IT has been struggling under the assumption that the whole (of quality of service management) is equal to the sum of all of the parts. In mathematics, this assumption is valid. In the management of IT services, it is totally false [Ref. 19]. Consider a simple example in which a typical transaction requires that each of the following pieces of equipment and software listed be available.

Component	Average Availability	Minutes of Downtime
LAN	99.97%	0.32
Local Server	99.95%	0.54
Building Hub	99.96%	0.43
Intranet router	99.99%	11.88
Remote Host	99.99%	1.08
Order Entry applic.	99.91%	9.72
Customer data base	99.92%	8.64
Inventory data base	99.96%	4.32
Average/Total	99.66%	36.94

Table 1. Availability. [Ref. 19]

Let's assume that the user group performing order entry demands 99.90% availability for the order entry "function." (Note the term function is being used to indicate the desired activity, such as entering a new order.) Looking at the table

above, most IT organizations will congratulate themselves on having met the client's requirements in all but one of the categories, and even that was very close to the objective. However, this is probably not correct. Unless there was the coincidence of all of the problems occurring at the same time, then the total time that the function was unavailable to the order entry department was much greater than the acceptable level. If the outages each occurred at a different time, then the total amount of time that the function will unavailable to the client was an average of 37 minutes. Also, the availability can be calculated by multiplying together the availability of each of the components. The result is 99.66% - far less than the target of 99.99%. In truth, there may be some overlap between the various events. This could reduce the total amount of time that the function was not available to the order entry clerks. However, overlaps could not conceivably reduce the minutes of outage to the target level.

From the perspective of the end user of the IT services, the availability of individual components simply does not reflect the quality of the service being provided [Ref. 19]. The only merit of this particular metric is it does reflect a concept the clients can relate to, even if it is from a different perspective. If data for individual components is all that is available, what is needed is a more sophisticated correlation and analysis of the data to translate it into meaningful information.

In terms of managing the quality of service to client organizations (end users), there is absolutely no value in reporting such arcane data points as packets dropped, network latency, repair interval for network outages, etc. With

a few rare exceptions, the clients simply do not understand their significance [Ref. 19]. Even if understood, these pieces of data are not measurements of the quality of the service being provided by the IT organization. It is true that such information may be valuable measurements for internal IT assessments. In fact, some groups develop what are termed "Technical Service Level Agreements." These are basically agreements that are used to look at how well the groups within IT are fulfilling their responsibilities.

We have established that technical measurements are valuable for special purpose assessments of quality, but in their raw form do not reflect the quality of service that is being received by end users. However, if device availability, as well as other detailed statistics, are not valid indicators of the quality of service being delivered to end users, what is? Before this question can be answered it is necessary to deal with the even more fundamental issue of defining what is meant by the term "quality of service."

C. DEFINING QUALITY OF SERVICE

First, it is necessary to assume the perspective of the client - the consumer of the service. The client's perspective is the only perspective that matters in this definition [Ref. 12]. This is true regardless of the industry. It is as valid in the fast food business as it is in the world of IT. The client is the one utilizing the service. The client's perception of the quality of the service being delivered will determine whether or not they are satisfied with that service and whether or not they will want to continue to work with that service provider.

There may be a temptation to consider the adoption of a client's perspective to be biased and possibly even one-sided. However, IT managers must remember that they are both providers and consumers of services [Ref. 12]. They are the clients when dealing with service providers such as the local telephone company. They are also consumers of services from their own organization. The IT department uses the services of various groups within IT to provide the total package of services that they offer to a client. For example, a typical IT department will have a group responsible for operating the corporate Intranet. Another group may be charged with the operation of the central host computers. Yet the IT department, or the CIO, ultimately uses these services together to provide an end-to-end service package for other departments in the business. In this mode, the IT department is a client for IT services [Ref. 12].

From the client perspective the quality of IT services can be thought of in fairly simple terms. The most important point to remember is that the client is concerned with business functionality. That is, they need to be able to use the IT service to perform a business function. They see the ability to enter customer orders as an IT service, not as a collection of independent services. They are not concerned with how the service is provided, or what components make the service possible [Ref. 12].

The concerns of the client can be classified into four broad categories [Ref. 12]. Those categories are listed below:

- **Availability**
- **Performance**

- **Accuracy**
- **Affordability**
- 1. **Availability**

Availability focuses upon the question whether the services are available to the client when the client wishes to use them. There is not a magic number of days or hours that provides the right answer for this category. Rather, the objective is to define in advance when the services will be required and meet that requirement. There is the most obvious meaning for this facet of quality of service - scheduled hours of operation. Even in this sense, it is important to note that it means actually being able to successfully perform the desired function. All of the required components must be available. Having the family car ready for your spouse to use for a trip to the local mall is an example of this. It is highly unlikely that your spouse will consider the car available for the trip when everything is in perfect running order, but one tire is flat.

There are also less obvious aspects of availability [Ref. 12]. Consider the case in which a company has a system that allows sales representatives to dial in using their laptop computer to check product availability while they are at a customer's office. In such a case, if a modem port is not available when a sales representative attempts to call, then for that sales representative the service is not available.

2. Performance

Words can hold vastly different meanings for different people. This is certainly true of the word performance. For IT personnel performance will

suggest such things as packets per second. However, for the client, the meaning is simpler - do the IT services function at an acceptable speed? The question of speed may be the question of response time in an on-line transaction system. In another case it might be the time that it takes to move a copy of a file from one office to another, or the time required to load an application to a desktop system from a server.

3. Accuracy

The third part of the user perspective of service is accuracy [Ref. 12]. This component is concerned with the question of whether or not the service performs accurately. An example of this is the question of whether or not e-mail is delivered to the correct recipient. Similarly, in the case of applying transactions to a database, it is essential that the change be applied to the proper version of the database. If services do not accurately perform their functions, high availability and high performance are worthless.

4. Affordability

While availability, performance and accuracy, are strictly related to the quality of service, the cost of the service cannot be ignored [Ref. 34]. There is a saying in application development, "Fast, cheap, good - two out of three can be obtained." This expression is true when it comes to delivering any service. There is always a tradeoff between cost of the service and the quality of the finished product, as well as the timeliness of the delivery. It may be possible for a service provider to raise the quality of the service that it is providing dramatically. However, the cost of such an improvement might be so great, that

it, if implemented, that it would cause the bankruptcy of the corporation. Therefore, while cost is not directly part of the quality of the service being delivered, it is a limiting factor, and cannot be ignored. Having loosely defined quality of service, the question of how to approach the management of that service and the level of the service that is provided will be discussed. The question of which specific indicators or measurements are appropriate to use will be deferred until later in this thesis. First we must look at how they will be used.

D. APPROACHING THE QOS PROBLEM

Managing the level of service that is provided is a little bit like the physics problem of creating a vacuum [Ref. 12]. Assume that a vacuum pump is attached to a sealed chamber and in a cycle lasting one hour, the pump is able to remove 50% of the air in the chamber. It is almost certain will never get to a total vacuum (that is, no air molecules left in the chamber), at least not in a reasonable amount of time. Likewise, service providers and their clients must accept the fact that they will never achieve perfect service, at least not at a cost that any business can afford. Perfect service, that is always available with excellent performance and 100% accuracy, is simply not possible [Ref. 12].

If perfect service is not possible, then it is necessary to determine what level of imperfection is acceptable. This determination is not something that can be done by the service provider alone, nor is it possible for the client to do this task. The service provider does not have adequate information about the value of the service to the client organization. The clients do understand their own need for the service and its value to the organization. However, the client does

not have a good understanding of the issues facing the service provider. The client may have enough technical expertise to realize that it is possible to substantially increase availability, but they will probably not know the financial and other implications of doing so.

The answer for determining the appropriate service level lies in approaching the problem as a partnership between the service provider and the client organization [Ref. 12]. (It must be remembered that although IT is most often thought of as the service provider, they are also major clients for many external organizations.) The ideal is for the two groups to come together to work out what could be best for the company, in a "win-win" approach to the problem.

Unfortunately, not all groups or individuals are willing or able to take such an approach. Some insist upon a win-lose approach to these discussions. Others come to the discussions with absolutely no ability or willingness to make concessions or adjustments of any kind. This latter case is best seen when a IT manager attempts to negotiate a service level agreement with the local telephone company. The IT manager might as well try to negotiate with the government's tax collectors. Even if one party is not a willing participant, it is important that the dialog takes place.

When discussions take place, everyone should be prepared to share as much factual information as possible about the costs and impacts of various service levels [Ref. 12]. This information can then make it possible to drive decisions based upon business impacts, rather than merely emotions.

What is the expected result from the dialogue regarding the quality of the service to be provided? Very simply, the dialogue should produce an understanding of what is desirable, possible and acceptable to all parties.

It must be remembered that when IT provides a service to a client or uses the services of an outside provider, it is still entering into a business relationship. It is easier for most people to recognize the existence of a formal business relationship when dealing with an external supplier. However, the business relationship is just as real when it is between two internal entities (i.e., IT and a client group) [Ref. 12].

Whenever a business relationship is established, expectations are set. In the case of internal IT services, the client expects that service will be delivered at an acceptable level of quality. IT funding may not directly depend upon the level of quality of the services that it provides. However, IT certainly has the right to expect that if it delivers the services requested at an agreed upon level of quality, the client will be satisfied with those services.

E. SERVICE LEVEL AGREEMENTS

What has been discussed in this chapter is a service level agreement (SLA). The term was discussed briefly in chapter one however has been avoided so far in this chapter because so many managers (both client and service provider) have had unfavorable experiences with them in the past and are reticent about considering them as viable options today [Ref. 15]. This raises two questions. First, why have SLAs been so notoriously unsuccessful in the past? Second, why should managers be willing to use them at this time?

1. Problems of the Past

In the past there have been numerous shortcomings in the process of creating and implementing SLAs [Ref. 15]. These shortcomings have, in many cases, led to the ineffectiveness of the SLA process and its subsequent abandonment. The shortcomings of the past can be divided into two very broad categories. The first encompasses those errors described earlier in this section. Clients and service providers have approached the process of developing an SLA as a contest - a test of the power of their respective organizations; the indicators chosen have not been meaningful, etc. Managers today are still susceptible to those problems. However, these are problems that have been identified. If two groups will agree to approach the SLA creation process using the guidelines in this document, the problems in this category can be minimized or avoided entirely [Ref. 15].

The second problem that dogged SLAs in the past was the lack of effective measurement tools [Ref. 15]. Much of the data available was either piecemeal or created manually. If the data was piecemeal, then its relevance to the SLA was usually marginal, or it forced the SLA to be written to accommodate what could be measured, even if it was not particularly meaningful. Data that was created manually was subject to another problem. Often it was not reliable due to human error, or due to deliberate manipulation of the input or the output (i.e. reports) to achieve the desired result. In some cases, individuals have been known to falsify reports. In recent years there has been a tremendous proliferation of hardware and software products (tools) for the monitoring and

control of networks and systems [Ref. 15]. These management tools have vastly increased the amount and the quality of the data available for use in SLAs. Furthermore, these advances in tool technology and product availability have also led to an increase in the reliability of the data that is generated. Finally, even when tools have been available to provide the means to capture the data required, they have provided a very narrow time perspective. That is, the tools do not generally serve as an archive of data. This has made it even more difficult to generate reports covering an appropriate span of time [Ref. 15].

2. Value of SLAs

The service level agreement can be a powerful tool for both service providers and their clients. As was noted previously in this chapter, the clients and the service providers have significantly different perspectives regarding the services being provided. It is, in fact, accurate to characterize these two communities as speaking different languages. These are languages that are not understood by the other group. This leads to frustration on the part of both groups. There are the repeated complaints such as "They don't understand." Or, "They just don't listen." The service level agreement offers a way to bridge this communications gap. By creating a document through process of negotiation, it should contain objectives that are realistic and, also, service level indicators that are meaningful. Furthermore, the ongoing process of administering the SLA requires regular meetings and dialog between the service provider and client organizations. These meetings foster improved communications. Over time, the increased amount of communication will help

elevate the quality of the communications that take place. Another interesting result of this process can be increased trust, as the two groups interact more frequently and communicate more successfully.

SLAs offer a particular advantage for service providers [Ref. 15]. They offer a tool that the service provider can use to manage client expectations. It forces the client to think about what it means to have good service, and how much they are willing to pay for it. Once documented, the SLA also insulates the service provider against the gradual escalation of client expectations that can occur in the case of undocumented service level commitments.

F. QOS SOLUTION

What characteristics does a robust QoS solution have? Consider the following [Ref. 16]:

- **End to End Policy Enforcement**—QoS is not something that can be applied in a single standalone appliance, hooked up to the network, and left in hopes that all the traffic behaves in an orderly manner. QoS must be applied end-to-end. Consequently, it must be platform, device, and media independent, operating at Layer 3 and above to ensure end-to-end functionality across multiple network devices (such as routers, switches, firewalls, access servers, and gateways) and link layers (for example, ATM, Frame Relay, or Ethernet). Cisco IOS™ software and technologies in Cisco devices share a related set of QoS tools or mechanisms that interoperate to enforce QoS from the network core to the very edges (See Figure 17).
- **Multiple Parameters**--Policies are based on how people use the network. Devices must have the flexibility to apply and enforce QoS based on several parameters that can closely reflect network managers defined policy. These parameters distinguish traffic flows based on IP or MAC address, application, user, time of day or location within the network. Administrators define policies using a combination of these parameters. For instance, a simple policy might read, "No Webcast traffic allowed across a certain 56K link

to Branch Office A." A more sophisticated policy might read, "Allow video conference traffic within a defined LAN segment on Mondays between 3 and 5 p.m. for users A, B, C, and D only. At all other times and for all other users, disallow video conference." Devices must have the intelligence to recognize applications and apply QoS in order to achieve business-driven policies. These devices include Layer 3-capable devices such as routers (for example, any Cisco router), Layer 2/3 devices such as routing-enabled switches (for example, the Catalyst 5500, 8500 switches).

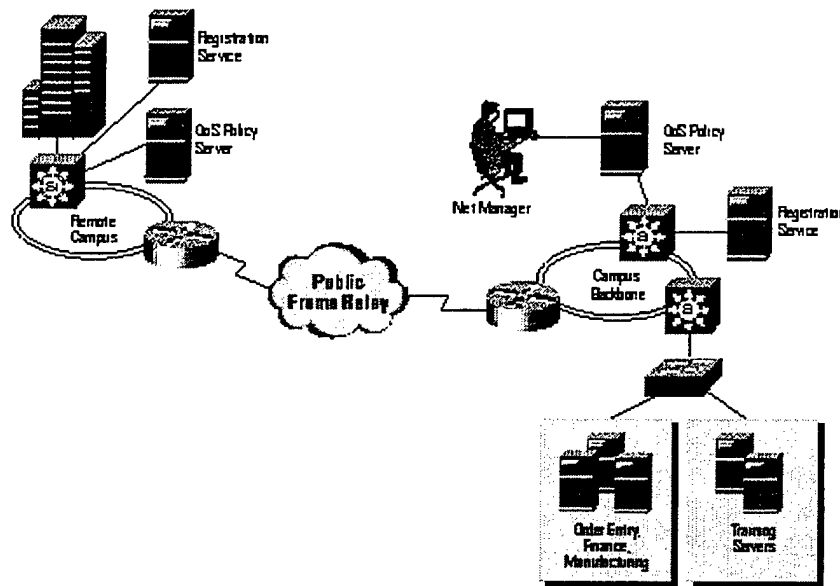


Figure 17. QoS Architecture. [Ref. 18]

- Classification--By definition, administrators need the ability to set different QoS levels for traffic as defined by the policy. In general, there are two types of QoS that assist in prioritizing traffic flows. They are as follows:
- Guaranteed Services reserve bandwidth for specific applications and/or users.
- Differentiated Services define a class of service for specific applications and/or users, e.g., high, medium, or low, where low could imply best effort, where there are no QoS controls; traffic is sent, as bandwidth becomes available.
- Centralized Control--Many software developers champion the ability of an application to signal QoS priority across the network. This assumes a "trusted" QoS implementation. With trusted QoS

controls, users are responsible for conforming to company QoS policy when they launch certain applications. The trusted QoS model may take control away from the network manager. The model may end up in the hands of people or applications that either may not have a full perspective about how the network must operate, or who do not understand or support the established business priorities that determine QoS policy [Ref. 10]. A further concern is whether to trust an end system to signal end-to-end QoS, which could include a WAN component. A possible alternate, "untrusted" approach is to centralize QoS implementation, removing the need for end-systems to police themselves and letting network managers set up and enforce the policy within the "trusted" network [Ref. 11]. Some networks may use a combination of trusted and untrusted QoS controls. As demonstrated in this chapter, network based policy enforcement more often results in a consistent policy deployment and enforcement. In this way, network managers can control network traffic and deliver mission-critical applications across the enterprise while still enabling traffic delivery for other applications.

- **Sophisticated QoS Tools**--Because there are so many different network elements and so many parameters required to successfully deploy and implement a QoS policy end to end, the associated sets of QoS tools are necessarily complex. They must be fully featured to enable network managers to build the intelligent networks they need.

G. INTERNET PROTOCOL QOS MECHANISMS

According to the Internet Engineering Task Force (IETF), there are two models for providing QoS [Ref. 13].

- **Integrated Services (int-serv):** The int-serv model is based on reservations-based traffic engineering assumptions. It reserves resources explicitly using a dynamic signalling protocol and employs admission control, packet classification, and intelligent scheduling to achieve the desired QoS.
- **Differentiated Services (diff-serv):** The diff-serv model is based on reservation-less traffic engineering assumptions. It classifies packets into a small number of service types and uses priority mechanisms to provide adequate QoS to the traffic. No explicit resource reservation or admission control is employed, although

network nodes do have to use intelligent queueing mechanisms to differentiate traffic.

1. **Integrated Services Model**

The following are the key characteristics of the int-serv model [Ref. 13].

(a) Flows

A flow is a stream of packets that originate from the same user activity e.g. a single application session. A flow may be identified by a variety of mechanisms; for example, IPv6 uses source address and flow label, and IPv4 uses source and destination address and destination port.

(b) Service Categories

int-serv defines the different service categories that are available, based on the delay and loss requirements, specifically:

- **Guaranteed Delay:** provides absolute guarantees on the delay and loss experienced by a flow (intended for non-adaptive real-time applications). Packets that conform to the reservation will not be lost, nor will they experience a delay exceeding the specified bound. Firm guarantees require a high level of resource reservation and may also result in worst-case guarantees that are significantly worse than the average case.
- **Controlled Load:** provides service equivalent to that of an unloaded network (intended for adaptive applications that perform well under lightly loaded network conditions). Most packets will not be lost, nor will they experience any queueing delay. No specific quantitative guarantees are provided.

(c) Traffic Specification

A flow's traffic is characterized by a TSpec (traffic specification).

The TSpec contains the following:

- **A Token-bucket specification:** A token rate **[r]** and a token bucket size **[b]**. These indicate the parameters for a token filter that can

be used to police the flow. Roughly, they are equivalent to the average data rate and maximum burst size for the flow.

- **A Peak data rate [p]:** An informational value that can be used to provide additional information for policing the flow and for providing appropriate bandwidth. However, its value may be set to the link rate or positive infinity, if no better value is known. A Minimum policed unit [m]: Used to allocate resources in the nodes as well as to compute the link overhead (by taking the ratio of link header to m).
- **A Maximum packet size [M]:** The largest packet that the application will send on this flow. This value should be no greater than the path MTU, since packets for flows with reserved resources are not fragmented within the network.

(d) Requested Service Specification

A guaranteed-delay flow can be further specified with an RSpec (requested service specification) to request a particular level of service. The RSpec contains the following:

- **Requested rate [R]:** Specifies the desired rate at which traffic is to be sent for the flow. It must be greater than r .
- **Slack term [s]:** Specifies the difference between the desired delay and the delay obtained by sending at rate R . It allows the network to adjust the allocated rate to meet the delay requirement (e.g., if it decreases the reserved rate, the queueing delay will increase, which reduces the available delay slack).

(e) Path Characterization

To interpret what QoS is available along a particular path, the resource reservation mechanism also advertises some path characteristics. These include:

- **int-serv capability:** Indicates whether there are any non-int-serv capable nodes along the path.
- **Number of int-serv hops:** Indicates the number of int-serv capable nodes along the path.

- **Available path bandwidth:** Indicates the minimum path bandwidth (not that for a particular flow).
- **Minimum path latency:** Indicates minimum delay for a packet, comprising propagation and processing delay, but not queueing delay.
- **Path MTU:** Indicates the maximum transmission unit along the path.
- **Delay parameters:** Indicates the cumulative delay along the path for a particular flow, comprising a rate-dependent term [c] and a rate-independent term [d].

(f) Resource Reservations

Resources must be reserved for flows in order to provide the requested QoS. This can be done via a dynamic reservation protocol, via manual configuration, or via a network management protocol. Int-serv is not tied to a specific mechanism and is deliberately defined to be independent of the actual mechanism used. It does, however, specify in a generic way what traffic and path characteristics are to be communicated [Ref. 13].

RSVP is a specific protocol designed to provide resource reservations. It is designed to work with int-serv but can be used with other service models as well. It has the following characteristics [Ref. 13].

- **Multicast environment:** RSVP was designed to work well in multicast environments, because the applications requiring resource reservations are multicast-oriented (e.g., conferencing). The protocol envisions simple communication between a set of senders and a set of receivers, with the actual communication path using (potentially overlapping) multicast trees.
- **Receiver-oriented:** To scale to large multicast environments, the protocol requires receivers to make reservations. Receivers request resource reservations based on the sender's traffic specification and path characteristics.

- **Receiver heterogeneity:** RSVP supports heterogeneous receivers by allowing each receiver to make its own reservation, perhaps distinct from others, even if the traffic is to be received from the same source. Intermediate nodes aggregate reservation requests.

(g) Soft State

The reservation state must be periodically refreshed by receivers, or it is timed out. This makes the protocol robust and well adapted to changing network conditions and changing reservation requirements. It also eliminates the need for reliable signalling mechanisms.

There are no in-built mechanisms for routing or packet scheduling: RSVP is just a signalling protocol. It relies on normal IP routing to compute the reservation route. Of course, it is desirable to include the resource request in the route lookup, but no standard QoS-aware routing protocols exist today for IP. RSVP is also not concerned with how nodes implement the reservation requests (e.g., admission control, packet classification, packet scheduling).

The following gives a brief overview of RSVP operations [Ref. 13].

- **Sender traffic specification:** The sender of traffic advertises its traffic using the Sender TSpec to construct an RSVP SENDER_TSPEC object, which is included in the RSVP PATH messages generated for the application. These PATH messages are sent towards the receivers (e.g., addressed to a multicast address). The SENDER_TSPEC object is not modified by the network and is delivered as-is to the receivers. The sending application also creates an initial path characteristics specification in an RSVP ADSPEC object, which is also carried in the PATH message. The ADSPEC object is modified by the network to reflect the path characteristics (e.g., bandwidth, MTU) as the PATH message propagates towards the receivers.

- **Receiver reservations:** Upon receiving the PATH message, the receiver uses the SENDER_TSPEC and ADSPEC to determine the traffic and path characteristics and then combines it with its own requirements to generate an RSVP FLOWSPEC object. This object contains the receiver traffic specification (TSpec) and, if needed, the requested level of service (RSpec). This is then transmitted in an RSVP RESERVE message that is routed along the exact opposite path that the PATH message traversed.
- **Merging reservations:** The network nodes are required to maintain state information for the PATH messages they have forwarded; they use this, along with the RESERVE message, to actually reserve resources and to route the RESERVE message back towards the source. In addition, nodes can merge requests from multiple receivers as the requests travel upstream.
- **Reservation Styles:** A receiver receiving traffic from multiple sources can choose to aggregate the traffic at intermediate nodes, using reservation styles. A reservation style can have explicit sender selection, in which a reservation is established for the senders explicitly listed in the reservation, or it can have wildcard sender selection, in which traffic from any sender is selected. Additionally, reservation styles can be distinct—a separate reservation is made for each sender—or shared—multiple senders can share the same reservation. The reservation style is specified using a filter spec. A receiver can specify the following filters:
 - **Wildcard-filter:** Uses wildcard sender selection and shared reservation. Traffic from all sources will be merged into a single flow reservation at the receiver.
 - **Fixed-filter:** Uses explicit sender selection and distinct reservations. Traffic from an explicitly identified set of sources is carried, using separate reservations for each sender.
 - **Shared-explicit-filter:** Uses explicit sender selection and shared reservations. Traffic from an explicitly identified set of sources is carried, using a single reservation for all senders. RSVP relies on regular IP routing for finding a path for the reservation set up. A number of experimental routing protocols are being considered for doing QoS-aware routing (e.g., QOSPF, I-PNNI).

- **Admission Control:** Int-serv relies on admission control to limit the amount of traffic that is admitted into the network, so that adequate resources are available to provide QoS to existing flows. Resource reservation requests are processed by nodes to determine if the new reservation can be accepted without adversely impacting existing flows.
- **Packet Classification and Scheduling:** Packets traversing a node need to be classified as belonging to a flow (or a class of flows) so that they can be serviced appropriately. This classification is based on some data being carried in the packet (e.g., in the IP or transport headers). In IPv4, the source and destination IP addresses, as well as transport port number, can be used. In IPv6, the source IP address and flow label can be used. Classified packets are assigned to some internal queueing function, which then services the queues based on QoS criteria. Intelligent queueing mechanisms must be used to deliver proper service.

2. Delay and Jitter

The delay experienced by the service traffic (packets) is an important aspect of the perceived quality of service. The aspects of delay have a different impact on different services [Ref. 11].

- End-to-end delay
- Delay variation or jitter

Interactive real-time applications (e.g., voice communication) are sensitive to end-to-end delay and jitter. Long delays reduce the interactivity of the communication.

Non-interactive real-time applications (e.g., one-way broadcast) are not sensitive to end-to-end delay but are affected by jitter. Jitter is usually accommodated by using a buffer at the receiver where received packets are stored and then “played back” at the appropriate time offset. The time offset (also called “playback point”) is determined by maximum jitter. Applications that

can adjust the playback point based on changes in the jitter value are called “adaptive” applications (e.g., vat). Packets that arrive after their playback point has passed are generally not useful to the application. Non-real-time applications are usually not delay-sensitive. However, because these applications may use delay measurements to control their traffic rate (e.g., TCP) or may have to buffer data until it is acknowledged (e.g., FTP), large or variable delays may affect the quality of these applications as well.

There are various components of end-to-end delay [Ref. 11].

- **Transmission delay:** the time it takes to put all the bits of a packet onto the link.
- **Propagation delay:** the time it takes for a bit to traverse a link (usually, at the speed of light).
- **Processing delay:** the time it takes to process a packet in a network element (e.g., routing it to the output port).
- **Queueing delay:** the time a packet must wait in a queue before it is scheduled for transmission.

At the endpoints, there may be additional delays in getting the packet from the network interface to the application and eventually to the user (e.g., delays in transferring the packet across the host bus, delays in copying the packet from kernel space to user space, delays in scheduling the application).

3. Throughput

The main aspect of throughput is the amount of bandwidth available to an application. This determines how much traffic the application can get across the network [Ref. 11]. Other important aspects are errors (generally related to link error rate) and losses (generally related to buffer capacity).

Certain applications can reduce their traffic rate in response to low throughput indications (e.g., reduce the fidelity of the encoding scheme). Such applications are called "rate-adaptive."

Throughput depends on the following factors [Ref. 11].

- **Link characteristics:** bandwidth, error rate
- **Node characteristics:** buffer capacity, processing power. As can be seen from the above discussion, a number of characteristics of various network elements, such as terminals/hosts, links, and switches/routers, determine what quality of service will be provided to applications, in terms of delay and throughput metrics.

4. Issues with the int-serv Model

Int-serv proposes a fundamentally new model for IP networks [Ref. 11]. It moves away from the best-effort model and introduces new services, such as guaranteed delay and controlled load. It requires explicit resource reservation, admission control, packet classification, and scheduling. It indicates how traffic and QoS requirements can be specified. There has been significant discussion on whether the int-serv model can scale to large backbones. Some of the difficulties are [Ref. 11].

- The number of individual flows in a backbone network can be very large. Maintaining state for all flows can require a lot of storage capacity.
- The number of control messages for resource reservations for a large number of flows can be large and may require a lot of processing power.
- Packet classification based on packet headers can be expensive on a high-speed data path.
- Security issues need to be resolved to ensure that unauthorized sources do not make spurious reservations.

- Policy issues need to be resolved to determine who can make reservations

The common wisdom is that int-serv is an appropriate model for a small intranet where there are a small number of flows and where security and policy issues can be managed easily [Ref. 10]. It may also be possible to deploy int-serv in larger networks (e.g., ISPs) for a limited number of flows (e.g., a few multicast sessions). Beyond that, large backbone networks will need more scalable mechanisms for differentiating traffic and providing differentiated services to them.

H. DIFFERENTIATED SERVICES MODEL EXPLORED

The continuing demand from ISPs for scalable and practical mechanisms to provide different levels of service to their users is primarily driven by the commercial interest in offering premium services for premium prices [Ref. 11]. It is also foreseen that, in the absence of some kind of service assurance, corporations will be highly unwilling to put mission-critical data on the Internet. The growth of such applications as voice-over-IP and virtual private networks is tied to the ability of the network infrastructure to provide some form of differentiated services for such applications [Ref. 11]. Due to the above-mentioned shortcomings of the int-serv model, a number of proposals were made to the IETF in the second half of 1997 to provide simpler, more scalable mechanisms for differentiated services. At this time, there is no agreement on the mechanism or even on the services that should be provided by this new model.

However, some of the key features of the various proposals that make them different from the int-serv model are [Ref. 47].

- **Coarser Granularity of Differentiation** - Backbone routers carry a large number of individual flows (where a flow is defined as a related stream of packets due to a single user activity such as TCP session). Maintaining reservation state per flow is intractable, because the Internet doubles in size every 6 months. Hence, the diff-serv proposals do not differentiate traffic per flow. Instead, there are a small number of well-defined classes which are provided different services based on delay and loss sensitivity.
- **No Packet Classification within the Network** - RSVP envisions that each router implements a packet classification function in order to provide different levels of service, which means that many fields of the packet header may need to be examined (e.g., for IPv4, the source and destination IP address, protocol, and source and destination port ids may need to be examined). This is a very expensive operation. Hence, in diff-serv, the proposal is to move packet classification to the edge of the network, where there are fewer packets to process. Edge routers classify and mark packets appropriately. Interior routers simply process packets based on their markings. **Note:** This implies that interior routers do not recognize individual flows, but that they deal with aggregate classes. This makes the solution more scalable.
- **Different Provisioning Models** - RSVP dictates dynamic resource reservation. This may result in a large number of control packets and it also requires dynamic admission control in each router. The diff-serv model moves admission control to the edge of the network. Further, it does not require dynamic reservations. Instead, it can use long-term, static provisioning to establish service agreements with the users of the network, whose traffic it can police at the ingress to the network. It can also use explicit traffic engineering to route packets from certain sources towards pre-determined paths, potentially bypassing congestion points in the network that may otherwise be on the path obtained by regular (non-QoS-aware) routing protocols.
- **No Absolute Service Guarantees** - The Guaranteed Delay service in int-serv provides hard bounds on the maximum delay that will be experienced by the packets. It does so by explicitly reserving resources along the path. In general, diff-serv will not provide hard guarantees unless a static path is provisioned specifically for a

particular flow. In diff-serv, interior routers do not distinguish individual flows, but, rather, deal with aggregated classes, so they cannot handle a particular flow preferentially.

The following description of diff-serv is based on the models proposed by [Ref. 11] and [Ref. 13]. The basic idea is to monitor the traffic that enters the network at the ingress node and check for compliance against some pre-defined service profiles. Based on this, packets can be marked as being "in" or "out" of their profiles. Inside the network, routers preferentially drop packets that are tagged as being "out." There are variations of the basic idea in terms of how many levels of precedence are defined for packets—two or more.

In addition to drop precedence, there can be some additional information in the packet headers that communicates the type of service (TOS) desired by the packet, particularly with respect to its delay sensitivity. For example, a premium service can offer the equivalent of a CBR connection (i.e., it is allocated according to its peak rate and is not oversubscribed, hence its bandwidth is always available and packets do not see any delay - this is akin to a "virtual link").

Another example can be an assured service that is characterized by bursty behavior and is provisioned using expected capacity, and hence its bandwidth is allocated statistically [Ref. 13]. Then, of course, there is the best effort service, which is expected to continue to form the bulk of the traffic. Additional levels of service may be defined. The service profiles are set up at the edge nodes based on customer subscriptions. Therefore, they are relatively static. The decision of whether to accept a new subscription can be made centrally, based on knowledge of network topology and capacity. Thus, a

network provider can provision its network according to the expected demand and subscriptions [Ref. 13].

Alternatively, it is possible to use a dynamic protocol to set up the profiles [Ref. 11]. In this case, the user would send in a service request, perhaps using an RSVP message. The edge node would direct the request to a bandwidth broker, which would make some decision and inform the edge node. There may be a need for bandwidth brokers to communicate among themselves (e.g., in a multi-network service set up). It is envisioned that in initial deployment, static provisioning would be used to size the network [Ref. 11].

To be able to provision the network appropriately, it is important to know not only how much traffic is entering the network, but also where it is going [Ref. 11]. Depending on the destination, resources may have to be provisioned on different paths. In the simplest case, the service profile indicates a desired resource between two specific end-points. The network provider has to ensure that the primary and secondary routes between the two points are adequately provisioned.

The problem becomes harder when there are multiple recipients (as in the case of a virtual private network) or when the recipients are not known beforehand. In the latter case, it may not be possible to provide any assurance of service at all, since the traffic might be taking any path in the network.

The main functional components of this architecture are as follows [Ref. 11].

- **Service Profiles:** A service profile indicates which traffic is to be treated differentially and what type of service is requested. The former may be indicated by setting a packet filter based on packet

header fields such as IP addresses. The latter can be specified using a token bucket filter (a rate and a burst size), along with some delay requirements. The service profile is likely to be set up administratively, although it is technically possible to use RSVP messages to set up profiles as well.

- **Packet Classification:** The ingress router must check all received packets against the service profiles it has to check if the packets should receive differential treatment. It may use packet filters to match packet headers and token filters to check conformance to the profile. Packets that do not meet profiles can either be discarded or sent into the network with higher drop precedence, depending on the nature of service to be provided. **Note:** The source can police and shape the traffic it is offering to the network in order to maximize the probability that the offered traffic will meet the service profile and receive the desired quality of service.
- **Packet Marking:** The ingress router must also mark the packets as they enter the network with appropriate values so interior routers can handle the packets differentially. The marking can use header fields. For example, for IPv4 packets, the TOS octet can be used. This has 3 bits for IP Precedence and 4 bits for Type of Service:

RFC 791 defines the values for IP Precedence bits as:

- 111 -- Network Control
- 110 -- Internetwork Control
- 101 -- CRITIC/ECP
- 100 -- Flash Override
- 011 -- Flash
- 010 -- Immediate
- 001 -- Priority
- 000 -- Routine

It has been proposed that one or more bits of this field be used to indicate drop precedence. If compatibility with RFC 791 is desired, higher values can indicate lower drop precedence; otherwise, it might be desirable to reverse the logic, since most IP packets today are sent with value 000.

RFC 1349 defines the following values for the 4-bit Type of Service field:

- 1000 -- Minimize delay
- 0100 -- Maximize throughput
- 0010 -- Maximize reliability
- 0001 -- Minimize monetary cost
- 0000 -- Normal service

Again, one or more bits of this field could be used to indicate the type of service being provided. For example, higher values could indicate higher delay sensitivity. At the other end, a single bit could be used to distinguish between best-effort traffic and delay-sensitive traffic.

Differential Queueing: In the interior routers, differentiated packets have to be handled differently. To do so, the router may employ multiple queues, along with some Class Based Queueing (CBQ) service discipline or simple priority queueing. Generally, delay-sensitive traffic will be serviced sooner, and loss-sensitive traffic will be given larger buffers. The loss behavior can also be controlled using various forms of Random Early Detection (RED). These disciplines use probabilistic methods to start dropping packets when certain queue thresholds are crossed, in order to increase the probability that higher-quality packets can be buffered at the expense of more dispensable packets. As an example, packets of different service types may be put into different queues, and, within a given service type, packets with higher drop precedence may be discarded earlier than those with lower drop precedence.

I. INTERACTION WITH MPLS

Another effort going on in the IETF is Multi-Protocol Label Switching (MPLS) [Ref. 21]. MPLS attempts to set up paths in a network along which packets that carry appropriate labels can be forwarded very efficiently (i.e., the forwarding engine would not look at the entire packet header, rather only at the label and use that to forward the packet). Not only does this allow packets to be forwarded more quickly, it allows the paths to be set up in a variety of ways: the path could represent the normal destination-based routing path, it could

represent a policy-based explicit route, or it could represent a reservation-based flow path.

Ingress routers classify incoming packets and wrap them in an MPLS header that carries the appropriate label for forwarding by the interior routers [Ref. 21]. In the MPLS model, the labels are distributed by a dynamic Label Distribution Protocol (LDP), which effectively sets up a Label Switched Path (LSP) along the Label Switched Routers (LSR). The LDP could be driven off destination-based routing (e.g., OSPF) or from reservation requests (e.g., RSVP) or some other policy-based explicit route. In some sense, LDP is creating label state in the network, but this is not so different from the normal forwarding tables created by routing protocols. It is important to note that this label state is not per packet or per flow, but usually represents some aggregate (e.g., between some source-destination pair). Therefore, the state produced by MPLS is manageable and scalable.

The MPLS model is compatible with the diff-serv model. The ingress routers can use service profiles to assign labels to packets. The LSPs could represent provisioned paths inside the network, and the labels carried in MPLS headers can be used to differentiate packets [Ref. 13].

The drop precedence of the packet can also be indicated in the MPLS header. It has also been proposed to use RSVP as the LDP so that a single protocol can be used for setting up various types of LSPs, including destination-based, policy-based, and reservation-based paths. There is no consensus regarding this [Ref. 21].

J. A NEW CHALLENGE: MANAGING QUALITY OF SERVICE (QOS)

QoS has been a critical requirement for wide-area networks for years. Bandwidth, delay, and delay variation requirements are at a premium in the wide area [Ref. 18]. The importance of end-to-end QoS is increasing because of the rapid growth of intranet and extranet applications that have placed increased demands on the entire network. QoS can protect mission critical applications from bandwidth hungry applications such as multimedia, web casting, and real-time video applications.

QoS provides a number of important roles [Ref. 18].

- Protects mission-critical applications such as ERP or sales automation systems.
- Prioritizes groups of users based on business functions such as sales and engineering.
- Enables multimedia applications such distance learning or desktop video-conferencing.

As seen in Figure 18, a typical policy networking architecture may have four components:

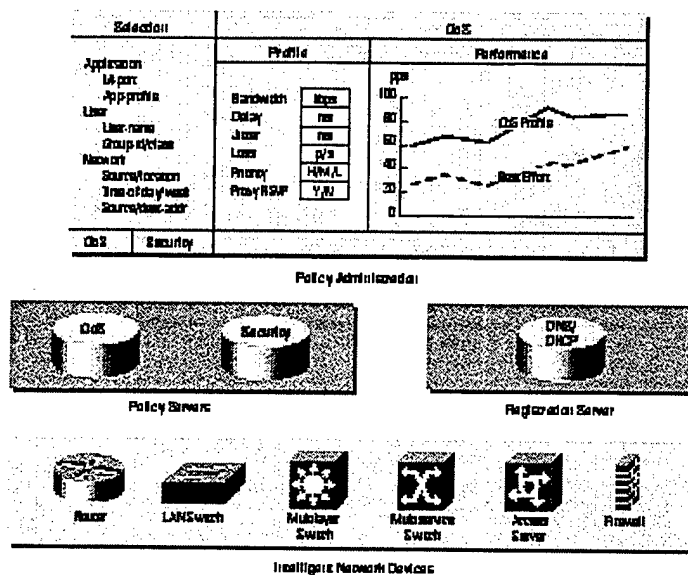


Figure 18. Policy Networking Architecture. [Ref. 12]

A network manager may need to provide different service levels for applications. For example, when a sales manager enters an order at the end of the quarter, the network can recognize the application and can prioritize it over other types of traffic [Ref. 12]. Today, the QoS mechanism in some software enable networks to control and predictably service a wide variety of networked applications and traffic types [Ref. 18]. Key QoS capabilities include router specific policies and signaling mechanisms as follows:

1. Queuing Techniques for Congestion Management on Outbound Traffic

A queuing technique can be set on a device's interface to manage how packets are queued to be sent through the interface. The technique chosen determines whether the traffic coloring characteristics of the packet are used or ignored [Ref. 18].

These queuing techniques are primarily used for managing traffic congestion on an interface, that is, they determine the priority in which to send packets when there is more data than can be sent immediately:

- First In, First Out (FIFO) Queuing
- Priority Queuing (PQ)
- Custom Queuing (CQ)
- Weighted Fair Queuing (WFQ)
- Weighted Round Robin (WRR)

a. First In, First Out (FIFO) Queuing: Basic Store and Forward

FIFO queuing is the basic queuing technique. In FIFO queuing, packets are queued on a first come, first served basis: if packet A arrives at the interface before packet B, packet A leaves the interface before packet B. This is true even if packet B has a higher IP precedence than packet A: FIFO queuing ignores packet characteristics.

FIFO queuing works well on uncongested high-capacity interfaces that have minimal delay, or when you do not want to differentiate services for packets traveling through the device.

The disadvantage with FIFO queuing is that when a station starts a file transfer, it can consume all the bandwidth of a link to the detriment of interactive sessions [Ref. 18]. The phenomenon is referred to as a *packet train* because one source sends a "train" of packets to its destination and packets from other stations get caught behind the train.

(1) Policy Requirements for FIFO Queuing Interfaces -

There are no specific requirements for creating policies on FIFO interfaces. Policies do not have to be defined on these interfaces. However, traffic shaping policies or traffic limiting policies on FIFO interfaces can set the rate limit on the bandwidth available to selected traffic. Traffic can be colored on a FIFO interface but it cannot be used to affect FIFO queuing [Ref. 18].

(2) FIFO's Relationship to Traffic Coloring - FIFO queuing

treats all packets the same: whichever packet gets to the interface first is the first to go through the interface. Traffic shaping and traffic limiting policy statements can affect the bandwidth available to a packet based on its color, but FIFO does not use the coloring value.

b. *Priority Queuing (PQ): Basic Traffic Prioritization*

Priority queuing is a rigid traffic prioritization scheme: if packet A has a higher priority than packet B, packet A always goes through the interface before packet B. When an interface's QoS property is defined as priority queuing, four queues may be automatically created on the interface: high, medium, normal, and low. Packets are placed in these queues based on previously defined priorities and policies. Unclassified packets are placed in the normal queue.

The disadvantage of priority queuing is that the higher queue is given absolute precedence over lower queues [Ref. 18]. For example, packets in the low queue are only sent when the high, medium, and normal queues are

completely empty. If a queue is always full, the lower-priority queues are never serviced. They fill up and packets are lost. Thus, one particular kind of network traffic can come to dominate a priority queuing interface [Ref. 18].

An effective use of priority queuing would be for placing time-critical but low-bandwidth traffic in the high queue. This ensures that this traffic is transmitted immediately, but because of the low-bandwidth requirement, lower queues are unlikely to be starved.

(1) Policy Requirements for Priority Queuing Interfaces - In order for packets to be classified on a priority queuing interface, policies must be created on that interface. These policies need to filter traffic into one of the four priority queues. Any traffic that is not filtered into a queue is placed in the normal queue. Traffic shaping policies or traffic limiting policies can be created to define an upper range on the bandwidth allocated to selected traffic [Ref. 18].

(2) Priority Queuing's Relationship to Traffic Coloring - Priority queuing interfaces do not automatically consider the IP precedence settings of a packet. If traffic coloring policies are created on inbound interfaces, and the coloring is desired to affect the priority queue, the priority queuing outbound interface must be created that recognizes the color value and places the packet in the desired queue [Ref. 18].

c. Custom Queuing (CQ): Advanced Traffic Prioritization

Custom queuing is a flexible traffic prioritization scheme that allocates a minimum bandwidth to specified types of traffic. Up to 16 of these custom queues can be created. For custom queue interfaces, the device

services the queues in a round-robin fashion, sending out packets from a queue until the byte count on the queue is met, then moving on to the next queue. This ensures that no queue gets starved, in comparison to priority queuing [Ref. 18].

The disadvantage of custom queuing is that, like priority queuing, policy statements must be created on the interface to classify the traffic to the queues. An effective use of custom queuing would be to guarantee bandwidth to a few critical applications to ensure reliable application performance.

(1) Policy Requirements for Custom Queuing Interfaces - In order for packets to be classified on a custom queuing interface, custom queuing policies on that interface must be created. These policies need to specify a ratio, or percentage, of the bandwidth on the interface that should be allocated to the queue for the filtered traffic. A queue can be as small as 5%, or as large as 95%, in increments of 5%. The total bandwidth allocation for all policy statements defined on a custom queuing interface cannot exceed 95% (QoS Policy System ensures 95% is not exceeded) [Ref. 18]. Any bandwidth not allocated by a specific policy statement is available to the traffic that does not satisfy the filters in the policy statements. QoS Policy System uses the ratio in these policies, along with the packet size specified when an interface is defined as a custom queue, to determine the byte count of each queue. The queues defined constitute a minimum bandwidth allocation for the specified flow. If more bandwidth is available on the interface due to a light load, a queue can use the extra bandwidth. This is handled dynamically by the device. All packets that have not been classified for custom queuing are placed in the default queue.

Traffic shaping policies can also be created or traffic limiting policies to define an upper range on the bandwidth allocated to selected traffic [Ref. 18]. Thus, the custom queue defines a minimum bandwidth, and the shaping policy or limiting policy defines an upper limit. When defining the bandwidth upper limit, the shaping or limiting policy must be executed before the custom queue policy, and it must filter the same traffic as the custom queue (or a subset of the same traffic).

(2) Custom Queuing's Relationship to Traffic Coloring -

Custom queuing interfaces do not automatically consider the IP precedence settings of a packet. If coloring policies are created on inbound interfaces, and the coloring to affect the custom queue is desired, a policy on the custom queuing outbound interface must be created that recognizes the color value and places the packet in the desired queue.

d. Weighted Fair Queuing (WFQ): Intelligent Traffic Prioritization

Weighted fair queuing acknowledges and uses a packet's priority without starving low-priority packets for bandwidth [Ref. 18]. Weighted fair queuing divides packets into two classes: interactive traffic is placed at the front of the queue to reduce response time; non-interactive traffic shares the remaining bandwidth proportionately. Because interactive traffic is typically low-bandwidth, its higher priority does not starve the remaining traffic. A complex algorithm is used to determine the amount of bandwidth assigned to each traffic flow. IP precedence is considered when making this determination. Weighted fair queuing is very efficient, and requires little configuration.

(1) Policy Requirements for Weighted Fair Queuing

Interfaces - Weighted fair queuing interfaces automatically create queues for each traffic flow. No specific policies are needed. However, traffic shaping policies or traffic limiting policies can be created to affect how select traffic is handled on the interface. A shaping policy or a limiting policy can control the bandwidth available to the selected traffic, whereas a coloring policy can change the relative importance of a packet and thus change how the interface queues the traffic [Ref. 18].

(2) Weighted Fair Queuing's Relationship to Traffic Coloring

- Weighted fair queuing is sensitive to the IP precedence settings in the packets. Weighted fair queuing automatically prioritizes the packets without the need for the creation of policies on the weighted fair queuing interfaces. However, if a coloring policy is created on the weighted fair queuing interface, it will affect how the selected traffic is queued. Weighted fair queueing can improve network performance without traffic coloring policies. When using coloring for weighted fair queuing, or WRED, realize that traffic starting devices can set precedence fields in their outgoing packets. Thus, it is important to set IP precedence on all network access points in order to avoid unintended traffic receiving high priority [Ref. 18].

e. *Weighted Round Robin (WRR): Traffic Taking Turns*

In Weighted Round Robin (WRR) there are four queues for each interface. Incoming traffic is placed in one of the four queues based upon

precedence. The system gathers IP precedence information from the service type field of the IP header. For an incoming IP packet, the first two (most significant) bits of the service type field determine the priority [Ref. 19].

(1) Modulo-N Hash - To select a next-hop from the list of N next-hops, the router performs a modulo- N hash over the packet header fields that identify a flow [Ref. 21]. This has the advantage of being fast, at the expense of $(N-1)/N$ of all flows changing paths whenever a next-hop is added or removed [Ref. 18].

(2) Hash-Threshold - The router first selects a key by performing a hash (e.g., modulo- K where K is large) over the packet header fields that identify the flow. The N next-hops have been assigned unique regions in the key space. By comparing the key against region boundaries the router can determine which region the key belongs to and thus which next-hop to use. This method has the advantage of only affecting flows near the region boundaries (or thresholds) when next-hops are added or removed. Hash-threshold's lookup can be done in software using a binary search yielding $O(\log N)$, or in hardware in parallel for $O(1)$. When a next-hop is added or removed, between $1/4$ and $1/2$ of all flows change paths [Ref. 19].

(3) Highest Random Weight (HRW) - The router uses a simple pseudo-random number function seeded with the packet header fields that identify a flow, as well as a next-hop identifier (address or index), to assign a weight to each of the N next-hops. The next-hop receiving the highest weight is chosen as the next-hop. This has the advantage of minimizing the number of

flows affected by a next-hop addition or deletion (only $1/N$ of them), but is approximately N times as expensive as a modulo- N hash. The applicability of these three alternatives depends on (at least) two factors: whether the forwarder maintains per-flow state, and how precious CPU is to a multipath forwarder [Ref. 19].

2. Traffic Shaping or Traffic Limiting Techniques for Controlling Bandwidth

Traffic shaping policies or traffic limiting policies can be created on a device's interface to manage how much of the interface's bandwidth should be allocated to a specific traffic flow [Ref. 13]. Policies can be based on a variety of traffic characteristics, including the type of traffic, its source, its destination, and its IP precedence settings (traffic coloring). Shaping differs from limiting in that shaping attempts to throttle traffic when it reaches the rate limits. The router buffers some of the traffic bursts. Only when the buffer fills are packets dropped. Whereas, limiting policies do not drop packets until rate limits are reached, then drop all packets that exceed the rate limit.

Unlike queuing techniques, which are part of an interface's characteristics, generic traffic shaping or traffic limiting is done through policies, which are defined in access control lists (ACLs). (Frame relay traffic shaping, FRTS, is defined in interface characteristics.) Queuing techniques only affect traffic when an interface is congested, or in the case of WRED, when traffic exceeds a certain threshold. With traffic shaping policies, flows are affected even during times of little congestion [Ref. 18].

These types of traffic shaping policies can be used [Ref. 13].

- Generic Traffic Shaping (GTS)
- Frame-Relay Traffic Shaping (FRTS)
- Limiting

a. Generic Traffic Shaping (GTS): Controlling Traffic on Non-Frame Relay Interfaces

Generic traffic shaping allows the setting of a bandwidth limit for specific types of traffic. For example, a policy that limits web traffic to 200 KB/sec can be created. This puts a cap on the bandwidth available to that traffic, ensuring that the remainder of the interfaces bandwidth is available to other kinds of traffic. In this example, if web traffic does not fill 200 KB/sec, other kinds of traffic can use the unused bandwidth [Ref. 13].

With generic traffic shaping, a buffer to accommodate traffic bursts can be defined, so that packets are not immediately dropped once the limit is reached. If a buffer is not defined, once the limit is reached, packets are dropped.

(1) Interface QoS Property Requirements for Generic Traffic Shaping - generic traffic shaping policies can be defined on any type of interface except those that use frame relay traffic shaping (FRTS). For custom queuing interfaces, a shaping policy can be created to form an upper limit for the bandwidth available to the selected traffic, and have the interface also apply a custom queuing policy to form a lower bandwidth limit for the traffic. However,

this is only available on selected combinations of IOS software versions and device models [Ref. 13].

b. Frame-Relay Traffic Shaping (FRTS): Controlling Traffic on Frame Relay Interfaces and Subinterfaces

Frame relay traffic shaping allows the specification of an average bandwidth size for frame relay virtual circuits (VC). The bandwidth allocated for bursty traffic can be defined, and control whether the circuit responds to notifications from the network that the circuit is becoming congested. By using FRTS, a minimum rate commitment can be defined for the virtual circuit, and accommodate the occasional need for greater bandwidth [Ref. 13].

Each virtual circuit (VC) is identified by a data link connection identifier (DLCI) and provides access to another endpoint of the frame relay (FR) cloud. The VC can be either a switched VC (SVC) or permanent VC (PVC). In the SVCs a connection establishment is made each time data needs to be sent over the network (like ATM) and negotiation of parameters occurs. For PVC, the connection is there all the time and its parameters are defined when it is ordered from the FR carrier.

If the FR is a fully meshed network, the FR cloud is viewed by the router like a shared media subnet, similar to an Ethernet interface in which few other routers connect to it. There may be a few VCs on the interface connecting to the FR network, but it will not be divided into logical subinterfaces.

It is very common that the FR interface is not fully meshed but is a collection of virtual point-to-point links [Ref. 13]. In this case subinterfaces will be defined on the interface connecting to the FR network and therefore only one VC

is defined on each of the subinterfaces. QoS Policy System applies FRTS definitions to all VCs defined on an interface or subinterface. Multiple VCs on a single interface or subinterface cannot be treated differently.

(1) Interface QoS Property Requirements for Frame-Relay Traffic Shaping - FIFO can be used, priority queuing, or custom queuing on frame relay subinterfaces. If priority queuing or custom queuing is used, policies on the interfaces or subinterfaces must be created that create the required queues. On a FR interface WRED can be used and weighted fair queuing (WFQ). By creating custom or priority queues, the automatic rate-limiting features of FRTS can be further modified. Parameters can be controlled through QoS are Rate (CIR), burst size (BC), excess burst size (BE) and adaptive rate (response to BECN notifications). Generic traffic shaping policies on an FRTS interface cannot be created. Traffic coloring policies on the interface can be created, however [Ref. 13].

c. *Limiting: Limiting Bandwidth and Optionally Coloring Traffic*

Limiting allows the setting of a bandwidth limit for specific types of traffic. For example, a policy can be created that limits web traffic to 200 KB/sec. This puts a cap on the bandwidth available to that traffic, ensuring that the remainder of the interface's bandwidth is available to other kinds of traffic. In this example, if web traffic does not fill 200 KB/sec, other kinds of traffic can use the unused bandwidth. Packets are dropped if traffic bursts exceed the limit [Ref. 13].

(1) Interface QoS Property Requirements for Rate Limited Traffic - limiting policies can be defined on any type of interface. For custom queuing interfaces, a limiting policy can be created to form an upper limit for the bandwidth available to the selected traffic, and have the interface also apply a custom queuing policy to form a lower bandwidth limit for the traffic. However, this is only available on selected combinations of IOS software versions and device models [Ref. 13].

(2) QoS Policy - Historically, configuration of QoS has been complex and error-prone due to its static nature, thus limiting its application to the wide-area network edge [Ref. 12]. QoS mechanisms had to be manually configured on each device. Now through Policy Networking, policy configuration is simplified to enable active policies in the network and deployment of QoS end-to-end. Policy Networking enables QoS policies based on application type, user group identity, and other classifications such as time of day, day of week, or even physical port information derived from the topology of the network itself [Ref. 19]. These identity classifications take advantage of policies in force and registration services as well as other important network information services like management systems that provide topology and device information. To set up a QoS policy, the network manager could use a "drag and drop" Policy Administration graphical-user-interface (GUI) to specify a policy based on business rules. A QoS policy binding is then created and activated by QoS Policy Servers (see Figure 19). The Common Open Policy Service (COPS) protocol provides policy exchange between the policy servers and the elected

policy software embedded in the intelligent network devices. The software will translates the policy binding into local QoS enforcement mechanisms such as Weighted Fair Queuing (WFQ) or Weighted Random Early Discard (WRED). After the policy is activated, specific policy-aware network devices identify and classify traffic and execute the appropriate policy dynamically without requiring manual intervention. As a result, the network manager can focus on specifying business policies and leveraging the intelligent network to recognize and enforce the policies automatically [Ref. 19].

QoS Policy Binding		
Identifier	QoS	Accept/Deny
SOLnet	High	Accept
NetMeeting	Medium	Accept
Video	<100 Kbps	Accept

Figure 19. QoS Policy Binding. [Ref. 21]

IV. NETWORK POLICY: SECURITY

A. Overview of Network Security

The role of a security policy is to ensure that each of the four fundamental components that make up computer security, Authentication, Access Control, Integrity and Confidentiality are adequately addressed. Typical questions that need to be answered when developing a network security policy are [Ref. 23]:

- What resources are we trying to protect ?
- Which people do we need to protect the resources from ?
- How likely are the threats ?
- How important is the resource ?
- What measures can be implemented to protect the resource ?
- How cost effectively and in what time frame can these be implemented ?
- Who authorizes users ?

These questions should be revisited periodically, as the network environment is very dynamic. The security policy identifies the threats that need to be protected against and defines the level of protection required. The security policy will itself contain several different policies, for example a Network Service Access Policy and System Specific Policies and will be based on a security strategy [Ref. 23].

An outsider can gain a complete understanding of an organization's business direction--its strategies and plans--through its internal communications.

Email, customer databases, supplier information, documents, spreadsheets, accounting files, customer orders, operation procedures, quality policies, the telephone system, if available to the public would tell anyone all they would need to know about a company. If in the wrong hands or compromised, it could be disastrous. The corporate jewels are indeed found in the bits. Ensuring the authenticity, accessibility, integrity, and confidentiality of corporate information should be a chief concern of those in the upper levels of management [Ref. 23].

B. A GROWING NEED FOR PROTECTION

Recent audit from the Government Accounting Office (GAO) evidence indicates that serious and widespread weaknesses in information security are jeopardizing government and civilian agencies ability to adequately protect the following: (1) assets from fraud and misuse; (2) sensitive information from inappropriate disclosure; and (3) critical operations, including some affecting public safety, from disruption. Significant information security weaknesses were reported in each of the 24 largest federal agencies, with inadequately restricted access to sensitive data being the most commonly cited problem [Ref. 23]. In a recent survey conducted by the Computer Security Institute in cooperation with the Federal Bureau of Investigation, 64 percent of the 520 respondents, which were from both the private and public sectors, reported computer security breaches within the last 12 months--a 16 percent increase in security breaches over those reported in a similar survey in 1997. While many of the survey respondents did not quantify their losses, those that did cited losses totaling \$136 million [Ref. 23]. In an October 1997 report entitled "Critical Foundations:

Protecting America's Infrastructures", the President's Commission on Critical Infrastructure Protection described the potentially damaging implications of poor information security from a national perspective, noting that computerized interaction within and among infrastructures has become so complex that it may be possible to do harm in ways that cannot yet be fully conceived.

In September 1996, GAO reported that a broad array of federal and private operations was at risk due to information security weaknesses and that a common underlying cause was inadequate security program management. In that report GAO recommended that the Office of Management and Budget (OMB) play a more proactive role in leading federal improvement efforts, in part through its role as chair of the Chief Information Officers (CIO) Council. Subsequently, in a February 1997 series of reports to the Congress, GAO designated information security as a new government wide high-risk area [Ref. 24]. More recently, in its March 31, 1998, report on the federal government's consolidated financial statements, GAO reported that widespread computer control deficiencies also contribute to problems in federal financial management because they diminish confidence in the reliability of financial management data [Ref. 25].

The need to safeguard information has been around since the beginning of modern time. Throughout history, certain types of information have been classified and disclosed only on a need-to-know basis. From secret passwords and battlefield messages to firewalls and cryptography, the technology may be

different but the purpose is the same. So why worry about information security?

Answer: Internet, Intranet, and Extranet applications [Ref. 23].

The Internet's greatest strength is also its greatest weakness: openness. The Internet is an excellent vehicle for providing cost-effective communications. Almost every business uses and benefits from the Internet's primary applications: electronic mail and the worldwide web [Ref. 26]. Taken a step further, the Internet can serve as a backbone for Intranet and Extranet applications including tele-conferencing. The Intranet is a private, internal network, typically based on web/browser-like user interfaces, that often uses the Internet for connecting remote users. It privatizes the Internet for internal communications. An Extranet uses the Intranet concept but instead of connecting internal users, it connects an organization's customers, suppliers, and business partners. All have become essential tools for business communications. All are inherently insecure [Ref. 26]. Securing these applications requires the creation and implementation of security policies that address the authenticity, accessibility, integrity, and confidentiality of the network's users and its information [Ref. 26]. (An excellent document to reference when setting computer security policies and procedures for Internet-based communications is RFC 2196 -- the Site Security Handbook.)

C. SECURITY STRATEGIES

1. Least Privilege

The principle of least privilege is to grant only those privileges that are required.

Systems that allow permission to be granted or revoked by operation providing fine-grain control are well suited to this [Ref. 26]. There is generally an overhead in terms of increased system maintenance. Adopting a least privilege strategy limits exposure to attacks and importantly limits the damage that can be done when an attack is successful. Many of the common security problems on the Internet can be viewed as failures to follow the principle of least privilege [Ref. 26].

2. Defence In Depth

The defence in depth strategy is summed up by the term "Belt and Braces", i.e. use as many security mechanisms as possible and arrange them so that they back each other up. One of the problems with firewall systems is that they provide an all or nothing solution to security. If the firewall is breached the internal network is a soft target. This was noted by Bellovin who coined the term "Hard on the outside, soft and chewy on the inside" [Ref. 26] to describe it. Some firewalls however do implement the principle of defence in depth using techniques such as Type Enforcement.

An important aspect of defence in depth that is often overlooked is the need to avoid common mode failures. For example if an attacker can exploit a security weakness in brand X's router then there is little point in having two of them. However, brand Y's router may not have the same weakness and therefore the principle of defence in depth is met. This principle is important in the context of firewalls as many of the products commercially available are variations of Trusted Information Systems Gauntlet or their tool kit.

3. Choke Point

A choke point is a single point through which all incoming and outgoing network traffic is funneled [Ref. 26]. As all traffic passes through a choke point it is the natural place to focus monitoring and control efforts such as Internet firewalls. It is also the natural place at which to break the connection with the external network if necessary.

Choke points are often criticized as an all-eggs-in-one-basket solution. This concern can be addressed by building some redundancy into the choke point. The key point is that the choke point provides control [Ref. 26].

The largest threat to a choke point strategy is if an attacker is able to bypass the choke point. As Firewalls generally act as choke points this is a significant issue, especially given the ease with which SLIP or PPP connections to Internet Service providers can be established.

As choke points can experience high levels of network traffic it is important to ensure that there is sufficient bandwidth available at the choke point to prevent a network traffic bottleneck. Any monitoring and logging software should also be able to cope with the level of network traffic [Ref. 26].

4. Fail Safe Stance

If a system is going to fail, it should be designed to fail into a safe state [Ref. 26]. This principle is particularly important in the design of Internet firewalls. Packet filters and application level gateways, both of which are discussed in the next chapter, should fail in such a way that traffic to and from the Internet is stopped.

5. Security Through Obscurity

This strategy is based on the hope that if you keep a low profile, would be attackers won't find you, and if they do, they will pass you by. Many companies do not publish the telephone numbers of their dial-in modems, only divulging the numbers on a need to know basis [Ref. 24]. Whilst this is a sensible precaution it is a poor basis for long term security. Information tends to leak out, and attackers are often skilled at eliciting information from staff using social engineering techniques.

Many organizations assume that an attacker won't be interested in them, and that they are therefore unlikely to be the target of an attack. The rationale behind this stance assumes that a site is targeted because the attacker is interested in the information stored on it [Ref. 24].

Such assumptions are, at best, naïve [Ref. 24]. Some attackers regard themselves as electronic freedom fighters battling against the commercialization of the Internet, others wish to sell the information they glean, others are motivated by the power and control they wield over the lives of the system administrators they affect, and many attacks are motivated by revenge.

Attacks generally involve several computers and a multitude of accounts. An attacker may capture accounts and gain unauthorized access to several systems before reaching his real target. A site can be compromised for no other reason than to provide a staging post for attacks on other sites, and to the attacker, it means little more than another IP address [Ref. 24].

6. Simplicity

Software is complex. As the size of a piece of software grows it becomes increasingly difficult to test all eventualities. Complex code will probably have unknown loopholes that an attacker can exploit. These loopholes may be convoluted but that will not prevent an attacker from trying to exploit them, some of the exploit attacks against sendmail have been extremely intricate.

Simplicity is an important factor in sound network defences. Application level gateway network Security systems should have all extraneous functionality removed and should be kept as small and simple as possible [Ref. 26].

7. Host Based Security

Host based security is probably the most common computer security model in current use [Ref. 26]. The major problem with the host based security model is that it does not scale well. The major impediment to effective host security in modern computing environments is the complexity and diversity of those environments. Even if all hosts are identical, the sheer number of them at some sites makes securing each of them difficult. Effectively implementing and maintaining host security takes a significant amount of time and effort, and is a complex task.

While the host security model might be appropriate for small sites, and whilst all sites should implement some level of host security, it is not cost effective for larger sites, requiring too many restrictions, and too many people [Ref. 26].

8. Network Based Security

Network security is designed to address the problems identified with host security. The network security model concentrates on controlling network access to hosts and services rather than on securing the hosts themselves [Ref. 27]. Network security approaches include building firewalls to protect trusted networks from untrusted networks, utilizing strong authentication techniques, and using encryption to protect the confidentiality and integrity of data as it passed across the network [Ref. 27].

D. SECURITY POLICY

RFC1244 - The Site Security Handbook presents a useful guide to developing a site security policy. It is currently being revised and is due to be re-issued shortly. The guide lists and discusses issues and factors that a site must consider when setting their own policies and makes some recommendations. Useful guidance on some of the higher level requirements necessary for network security policy to be effective can be found in [NIST95].

To be effective, policy requires visibility [Ref. 28]. Visibility aids implementation of policy by helping to ensure policy is fully communicated throughout the organization. Management presentations, videos, panel discussions, guest speakers, question/answer forums, and newsletters increase visibility. The organization's computer security training and awareness program can effectively notify users of new policies. It also can be used to familiarize new employees with the organization's policies.

Computer security policies should be introduced in a manner that ensures that management's unqualified support is clear, especially in environments where employees feel inundated with policies, directives, guidelines, and procedures [Ref. 28]. The organization's policy is the vehicle for emphasizing management's commitment to computer security and making clear their expectations for employee performance, behavior, and accountability.

To be effective, policy should be consistent with other existing directives, laws, organizational culture, guidelines, procedures, and the organization's overall mission [Ref. 28]. It should also be integrated into and consistent with other organizational policies (e.g., personnel policies). One way to help ensure this is to co-ordinate policies during development with other organizational offices.

1. Site Security Policy

The Site Security Policy is an overall policy regarding the protection of the organization's information resources [Ref. 30]. This includes everything from document shredders to virus scanners, and remote access to floppy disk tracking. At the highest level, the overall organizational policy might state:

Information is vital to the economic well-being of the organization. Every cost-effective effort will be made to ensure the confidentiality, integrity, authenticity, availability and utility of information. Protecting the confidentiality, integrity, and availability of information resources is a priority for all employees at all levels of the company [Ref. 30].

Below this come site-specific policies covering physical access to the property, general access to information systems, and specific access to services

on those systems. The firewall's network service-access policy is formulated at this level.

2. Network Service Access Policy

The Network Service Access Policy is a higher-level, issue-specific policy which defines those services that will be allowed or explicitly denied from the restricted network, plus the way in which these services will be used, and the conditions for exceptions to this policy [Ref. 31].

While focusing on the restriction and use of internetwork services, the network service access policy should also include all other outside network access such as dial-in and SLIP/PPP connections. This is important because restrictions on one network service access can lead users to try others. For example, if restricting access to the Internet via a gateway prevents Web browsing, users are likely to create dial-up PPP connections in order to obtain this service. Since these are non-sanctioned, ad hoc connections, they are likely to be improperly secured while at the same time opening the network to attack.

For a firewall to be successful, the network service access policy should be drafted before the firewall is implemented [Ref. 31]. The policy must be realistic and sound. A realistic policy is one that provides a balance between protecting the network from known risks while still providing users reasonable access to network resources. If a firewall system denies or restricts services, it usually requires the strength of the network service access policy to prevent the firewall's access controls from being modified or circumvented on an ad hoc basis. Only a sound, management-backed policy can provide this defence

against internal resistance. Here are the typical network service access policies that a firewall implements:

Allow no access to a site from the Internet, but allow access from the site to the Internet; or, in contrast,

Allow some access from the Internet, but only to selected systems such as information servers and e-mail servers [Ref. 31].

Firewalls often implement network service-access policies that allow some users access from the Internet to selected internal hosts. This access should be granted only if necessary and only if it could be combined with advanced authentication [Ref. 31].

3. Firewall Design Policy

The Firewall Design Policy is a lower-level policy which describes how the firewall will actually go about restricting the access and filtering the services as defined in the network service access policy [Ref. 32].

The firewall design policy is specific to the firewall. It defines the rules used to implement the network service access policy. This policy must be designed in relation to, and with full awareness of, issues such as firewall capabilities and limitations, and the threats and vulnerabilities associated with TCP/IP. Firewalls generally implement one of two basic design policies:

Permit any service unless it is expressly denied;
or
Deny any service unless it is expressly permitted [Ref. 32].

A firewall that implements the first policy allows all services to pass into the site by default, with the exception of those services that the network service access policy has identified as disallowed. A firewall that implements the second

policy denies all services by default, but then passes those services that have been identified as allowed. This second policy follows the classic access model used in all areas of information security [Ref. 32].

The first policy is less desirable, since it offers more avenues for getting around the firewall. For example, users could access new services currently not denied by the policy (or even addressed by the policy). For example, they could run denied services at non-standard TCP/UDP ports that are not specifically denied by the policy. However, certain services, such as X Windows, FTP, Archie, and RPC are difficult to filter. For this reason, they may be better accommodated by a firewall that implements the first policy. Also, while the second policy is stronger and safer, it is more restrictive for users; services such as those just mentioned may have to be blocked or heavily curtailed. Certain firewalls can implement either design policy but one particular design, the dual-homed gateway, is inherently a "deny all" firewall. Systems which require services that should not be passed through the firewall could be located on screened subnets separate from other site systems.

In other words, depending on security and flexibility requirements, certain types of firewalls are more appropriate than others, making it extremely important that policy is considered before implementing a firewall [Ref. 32]. Failure to do so could result in the firewall failing to meet expectations. Specifics on firewall design are discussed later in this chapter.

4. System Specific Policies

System-specific policy is often implemented through the use of access

controls. For example, it may be a policy decision that only two individuals in an organization are authorized to run a particular program. Access controls are used by the system to implement (or enforce) this policy [Ref. 33].

5. Incident Handling

When a site that is not protected comes under sustained attack one of two things can happen. The site can rapidly develop a policy and defences or it can withdraw from the Internet [Ref. 31]. Internet security incidents, such as break-ins and service disruptions, have caused significant harm to several organizations' computing capabilities. Many organizations have an ad hoc response when initially confronted with an attack which can exacerbate the damage caused by the attack. For this reason it is often cost-effective to develop an in-house capability for the quick discovery of, and controlled response to, network security incidents.

The primary benefits of an incident handling capability are the ability to contain and repair damage resulting from network attacks. An incident handling capability also assists an organization to prevent, or at least to minimize, damage from future incidents. Incidents can be studied internally to gain a better understanding of the organization's vulnerabilities so that more effective safeguards can be implemented [Ref. 31].

6. Disaster Recovery

It is prudent to assume that an attack may fundamentally compromise an organization, for example deleting large amounts of data. It is for such eventualities that organizations develop disaster recovery plans. The basic steps

in establishing a disaster recovery plan are [Ref. 33]:

- Identify the mission or business critical functions.
- Identify the resources that support the critical functions.
- Anticipate potential contingencies or disasters.
- Select contingency planning strategies.
- Implement the contingency strategies.
- Test and revise the strategy.

E. THE KEY ELEMENTS OF SECURITY

1. Authenticity

Authenticity is the way in which a user or device proves its identity. The simplest form of user authentication is a user ID and password. More secure methods involve one-time passwords or some form of two-factor identification: something you know and something you possess [Ref. 29].

Anyone with a bank ATM card uses this type of authentication: your memorized PIN and the actual card. More elaborate systems include electronic tokens based on some type of time synchronization scheme or some type of challenge and response method. Extremely sophisticated schemes are based on biometrics in which a part or characteristic of one's body, fingerprint, hand, retina, face, voice is used to identify an individual. Device authentication is an electronic method of allowing a specific device to participate on a network. This is done by assigning devices to a specific switch port; allowing predetermined MAC addresses; restricted access to specific network addresses (IP address) or network protocol types (IP, IPX, etc.); or some combination of all four [Ref. 29].

2 Access Control

Access control is an extension of authenticity. It determines who or what has access to which network object or service. This is the category in which firewalls are placed. Through network addressing schemes, specific devices, applications, protocols, users or user groups (subnetworks), and regulate rights and privileges can be identified. Implementation of these is subjective. Ideally, they are implemented according to well-defined security policies that coincide with the nature of the business, the industry, company goals and objectives, and the culture and work environment [Ref. 29].

3 Integrity

Integrity deals with the condition of the received data and how it relates to the original sent data. It is essential that the sender and the receiver are working with the exact same data. Data packets can be captured while in route, manipulated by adding, removing, or changing data, and sent on to the target destination as if it came directly from the originator. Digital signatures and message integrity checks are technologies used to provide assurance that a messages arrives as intended [Ref. 29].

4. Confidentiality

Confidentiality is how privacy of information is guaranteed. The bulk of our Internet communication; e.g., email, is sent in clear text, similar to a typical postcard. Anyone who so desires has the ability to read what is sent. We accept this without much concern. However, some information is more valuable. We would rather conceal its identity. We do so by placing it in envelopes. Data

envelopes involve the actual "scrambling" of the data. The content can be unscrambled only by those who have the unscramble keys. Encryption is the area within the data security realm that handles this scrambling/unscrambling. Encryption is the foundation of a virtual private data network (VPN): a private network created over a public one [Ref. 29].

5. Security Must be Implemented in the Network

Many security policies can be implemented on switches. Switches are the foundation of today's modern networks [Ref. 30]. All data running over the network traverses the switches. End-user commands, server communications, other shared device communications, all occur over or through switches. Switches are ideal platforms for implementing access, firewall, authentication, authorization, and logging security services [Ref. 30]. Many security policies are ideally suited to be deployed throughout the network--into the workgroups, across the backbone, and at both the network ingress and egress points. As companies begin to move away from shared media-based networks, and toward switched-based networks with integrated layer-two and layer-three capabilities, the switch becomes a ubiquitous device. Because security is an organizational issue, implementing it on a broad basis requires security-savvy network devices. And switches, because of their abundance within most networks, are ideal places for security implementations [Ref. 30].

6. Authentication Services

Most organizations employ re-useable password schemes: a user ID and password used over and over. Although they are convenient, they are very

insecure. Because most passwords are sent from the client to the host/server in clear text (not encrypted), stealing passwords directly off of the wire are easy [Ref. 28]. Additionally, most passwords are easily guessed. Users may find they will use more than one password throughout a typical workday. Most network operating systems and administrators force users to change their passwords every month or so. To simplify their life, they often write their passwords down in an easily accessible location: on a Post-It note placed under the keyboard or phone, inside a daily planner or calendar, or on a whiteboard. Not the ideal way to ensure network security. There are better methods offering greater security and easier use. One-time password schemes are proving to be the better methods [Ref. 28]. They typically leverage more than one authentication factor. These factors include:

- Something you are (human uniqueness)
- Something you know (user ID, password, PIN)
- Something you possess (tokens and smart cards)

The most popular and affordable schemes are two-factor: something you know and something you possess. Authentication schemes that involve multiple factors prove identity with greater confidence than single factors. Whatever method is used, the process must be fast, secure, and easy to use. Popular schemes include [Ref. 28]:

- S/Key (an Request For Change (RFC)-based one-time password scheme)
- Time-Based (synchronized client/server system; i.e., SecurID)

- Challenge & response (client/server system leveraging portable devices; i.e., Defender)
- Smart cards (intelligent devices; i.e., card readers)
- Kerberos (client/server software system based on a ticketing mechanism)
- Fortezza (NSA-defined credit-card-size security device that authenticates and encrypts)
- Remote Access Dial In User Service (RFC-based using challenge/response options) [Ref. 23].

a. *Secure Access to Layer-Two Groups*

A growing trend in authentication is to segment the network into privileged zones or groups. A user, once authenticated, gains access to one or more pre-authorized user groups. Switch, router, and/or firewall rules regulate communications between groups. This method provides a higher level of security than traditional network sign-ons. Instead of signing on to the entire network, the authenticated user only gains access to one or more layer-two broadcast domains [Ref. 23].

Historically, access to these layer-two domains has been tied to a single characteristic of the user's device: a switch port, Media/Medium Access Control (MAC) address, protocol type, or network address [Ref. 23]. Additional security can be realized if multiple characteristics are used (device authentication). In addition, some organizations have mobile users within a campus environment. Tying a device to a user is not always possible. Instead of granting layer-two authorization based on the user's device, there's a way now to base layer-two access on whom the user is (user authentication) [Ref. 23].

This user authentication service leverages modern authentication techniques--one-time password schemes like S/key, time synchronization, challenge and response, and Remote Authentication Dial-In User Service (RADIUS). RADIUS is a protocol by which users can have access to secure networks through a centrally managed server. RADIUS provides authentication for a variety of services, such as login, dialback, SLIP, and PPP [Ref. 24].

In a RADIUS query, Media Access Exchange (MAX) provides a user ID and password to the server. The server sends back a complete profile, which specifies routing, packet filtering, destination-specific static routes, and usage restrictions specific to the user. In addition, the MAX can use the data in the RADIUS database to create and advertise static routes and to place outbound calls. MAX is a system-level network access unit, with a cage and backplane into which Multiband or Pipeline cards can be inserted to configure it for various application requirements. It supports up to 32 host ports or direct Ethernet connection and up to 8 Mbit/s to the network. It supports multiple applications, including remote LAN access, leased line backup and individual video-conferencing units, as well as connecting video-conference Multiple Chip Units (MCUs) to the digital dial-up network.

User Datagram Protocol/Internet Protocol (UDP/IP) provides the communications channel between a RADIUS client and server, with messages acknowledged. The primary advantage in using RADIUS to authenticate incoming calls is that you can maintain all user information off-line on a separate UNIX-based server. You store virtually all Connection Profile information on the

RADIUS server in a flat American Standard Code (ASCII) database. This server can accept authentication requests from many machines, which makes swapping out one dial-in network server for another much easier. (For more information, refer to RFC 2058 and 2059.) These services interoperate with existing IP management schemes like Dynamic Host Configuration Protocol (DHCP) [Ref. 24].

b. Device Authentication

Device authentication is the binding of a switch port to specific MAC address and/or a network protocol or network address. This mechanism is ideal for non-mobile systems, like printers, servers and certain types of workstations [Ref. 26]. If the device moves from its allocated port location, it will not be allowed to communicate over the network. If the network address or protocol does not match the port number and MAC address, it will not be authorized to participate on the network. This is an anti-spoofing mechanism.

An example is found by contrasting security of data center devices to non-data center devices. Data centers are physically secured, often with "keyed" entries. Work areas are typically less secured than data centers. Because more people come and go in the non-data center work area, additional security may be required on some devices. Device authentication provides an added level of security by electronically locking a device to a specific port. Authenticated ports can co-exist on the same switch as non-authenticated ports. See Figure 20.

c. User Authentication

User authentication is the process or service offered by the switch and an associated authentication management server that gives users authorization to participate in one or more layer-two groups or Virtual Local Area

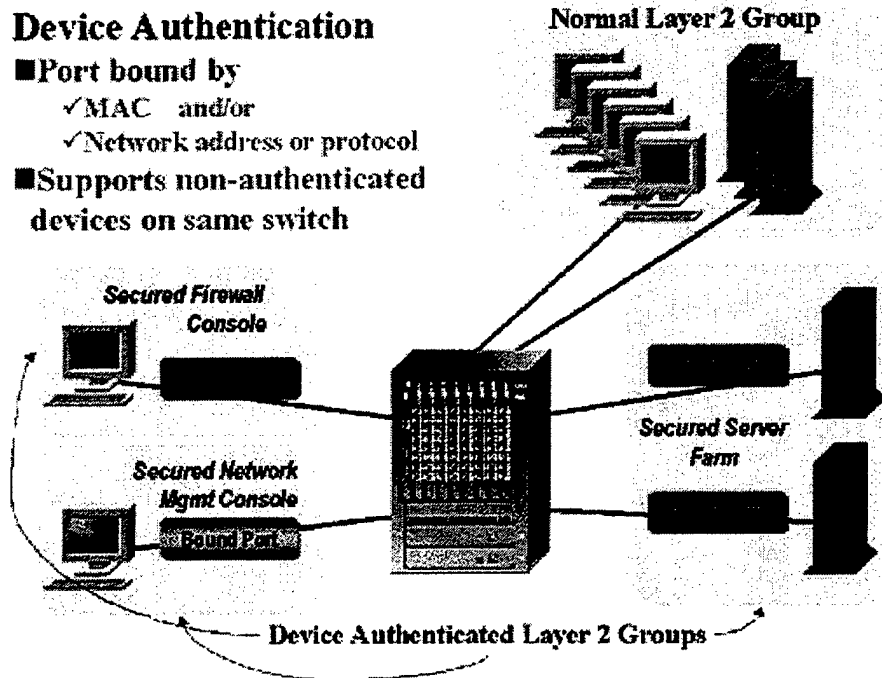


Figure 20. Device Authentication. [Ref. 28]

Networks (VLANs). This mechanism ties layer-two access to the individual user, not to the user's device. This couples the need for secure access to the need to accommodate user mobility. See Figure 21. There are three elements to the user authentication process [Ref. 28]:

- Authentication server
- Authentication agent
- Authentication client

User Authentication Architecture

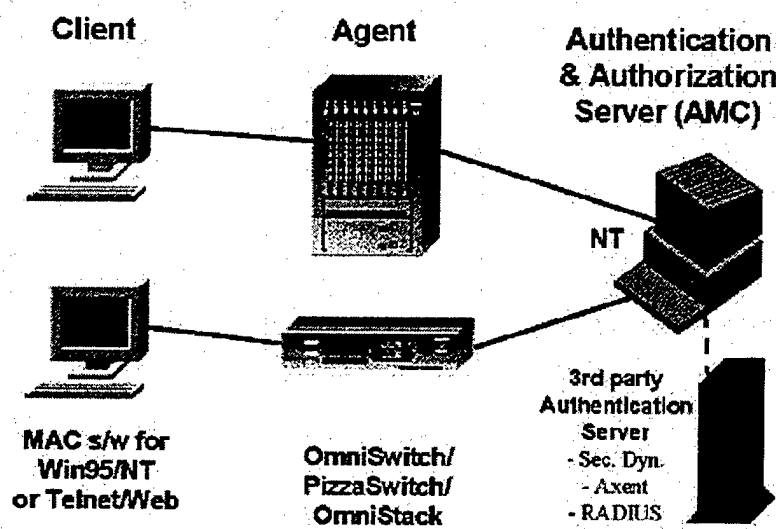


Figure 21. User Authentication. [Ref. 28]

1. Authentication Server. Known as the Authentication Management Console (AMC) is a software application from Check Point Corporation that is supported on Windows NT (and soon on UNIX platforms). At least one authentication server is required per network. A network administrator uses the server to maintain the information needed for user authentication. The authentication server provides a graphical interface that enables network administrators to manage the information for each user or user group [Ref. 28]. The AMC maintains a secure communication path with the agent in the switch (based on S/key). The authentication agent sends the user ID, physical location (switch port), and challenge/response of the user attempting to log in to the server over this communication path. The server sends the appropriate login

challenge and the VLAN authorization information to the authentication agent. The AMC uses the user ID response to the challenge and the time-of-day to determine if the user is allowed to connect to the network. The AMC maintains a log with the following information on every authentication attempt [Ref. 28]. The authentication server supports the following native authentication methods.

- S/key
- Firewall password
- Operating system password

If a third-party authentication method is used, the authentication server maintains an additional connection to the third-party server. The challenge and validation necessary to perform authentication is generated by the third-party server and sent transparently by the authentication server to the user attempting to log in. The authentication server supports the following third-party authentication mechanisms [Ref. 28].

- Security Dynamics' SecureID
- Axent's Defender
- RADIUS

All user authentication and network service access information is located on the authentication or third party server. This information is contained in a variety of formats suitable to the customer's requirements. The switches operate as clients to the authentication(s). The client sends authentication requests to the server and acts on responses sent back by the server [Ref. 28].

2. Authentication Agent. It is a software loadable module that resides on the switch. The authentication agent communicates with the authentication client, the authentication server, and special switch software known as AutoTracker to support the authentication process [Ref. 28]. The agent supports a MAC-based flow and a TELNET-based flow to the authentication client. The client sends the user ID, password, and any additional challenge response over this flow. The agent supplies the login prompt and any additional challenge supplied by the authentication server to the client over this interface. The agent also conveys the successful or unsuccessful results of the login process for display by the client. When the agent is initialized, it establishes an authenticated Transmission Control Protocol (TCP) session with the AMC. When the agent receives the list of authorized group names for the user, the agent forwards the list to the switch's AutoTracker management software. AutoTracker uses this list to authorize the Group/VLANs for the source MAC address into the filtering database. When the authenticated group is configured, the administrator can define the group as being the default for the protocol family that is assigned to the group; e.g., IP, IPX, AppleTalk or DECnet.

3. Authentication Client. It is any end-user device connected to a switch's switch port. The switch must be able to see the end-user device's MAC address. Hence, there cannot be a router or other gateway between the end user and the switch port. The user, either via MAC-client software file for Windows 95/NT end stations or via TELNET, issues a login request. The switch's agent software forwards the request to the authentication

server. If successful, the user device's MAC address is moved into the group or groups that the user is authorized to participate in. The user can log out, move to another device, log on again, and those same privileges will be granted. The only difference will be that there will be a different MAC address in the switch's registry [Ref. 28].

F. FIREWALL SERVICES

Much of the literature on firewalls concentrates on diagramming the numerous possible configurations of routers, host systems, interfaces, and subnets. It is imperative, however, not to lose sight of the broad definition of a firewall as a part of security policy [Ref. 16].

A firewall is a component or set of components that restrict access between a protected network and the Internet, or between other sets of networks [Ref. 27]. The typical application is to regulate the ingress to the private network from the public one (Internet). Intranet and extranet applications raise new issues traditional security measures do not address--protection from "trusted" or "pseudo-trusted" users.

A firewall **can** do the following [Ref. 27]:

- Be a focal point for security decisions (choke point)
- Enforce security policy (traffic cop)
- Log activity (misuse record)
- Limit your exposure (keep intrusions from spreading)

A firewall **cannot** do the following [Ref. 27]:

- Protect against every malicious insider (many yes, not all)

- Protect against connections that do not go through it (some go around it)
- Protect against completely new threats (known threats only)
- Protect against viruses (though they can work with products which do [Ref. 27])

1. **Types of Firewalls**

There are four classes of firewalls [Ref. 28]:

- **Packet filters**
- **Application gateways**
- **Circuit-level gateways**
- **Stateful inspection engines**

a. Packet Filtering Firewalls

Allow or block packets based on IP source and destination addresses, protocol (TCP, UDP), and TCP and/or UDP source/destination ports. Because this capability is included as part of a router's base software, this is the most widely implemented firewall type. These products offer decent performance and good scalability but have limitations on their depth of security and manageability. They deal with well-known ports, are susceptible to spoofing and common attacks, and use an error-prone implementation process [Ref. 28].

b. Application Gateways

Also known as bastion hosts or proxy servers; use special programs written for specific services (TELNET, ftp, http, etc.) to forward or deny packets based on security policies. The programs, or proxies, are implemented

on a host which typically has a hardened operating system (except for NT machines, but that's a different story) in order to remove bugs and holes. Instead of the user and applications talking directly with each other, they communicate through the specific proxy for that service request. They are very secure and easy to manage but have poor performance and do not scale. They perform poorly because the proxy process has tremendous latency: packet goes to the OS, then to the proxy, and back to the OS, with processing done at each step [Ref. 28].

c. *Circuit-level Gateways*

Relay TCP/UDP connections. They do not run proxies. Instead, users connect to TCP ports on the gateway that connects to a destination on the other side. The gateway program(s) copy bytes back and forth. The gateway acts as a wire [Ref. 27]. Circuit gateways are fast. However, they are troublesome to implement and require client software. Uncooperative insiders through the advertisement of unauthorized services can compromise them.

d. *Stateful Inspection*

It is a firewall technology invented by Check Point [Ref. 28]. Its inspection module analyzes all packet communication layers, and extracts the relevant communication and application state information. The inspection module understands and can learn any protocol and application. It resides in the operating system kernel, below the network layer. By inspecting communications at this level, it intercepts and analyzes all packets before they reach the operating systems. No packet is processed by any of the higher protocol layers

unless it verifies that it complies with the enterprise security policy. The inspection module has access to the "raw message," and can examine data from all packet layers. It analyzes state information from previous communications and other applications. It examines IP addresses, port numbers, and any other information required determining whether packets comply with the enterprise security policy. It stores and updates state and context information in dynamic connection tables [Ref. 28]. The advantages of stateful inspection are speed, ease of management, security depth, and scalability.

2. Switch-based Firewalls

Switches offer an attractive platform to deploy a firewall solution for many external applications and most internal applications [Ref. 29]. The strength of a switch-based firewall implementation lies in the ability to leverage an installed networking device, its higher port count, built-in redundancy, higher performance potential, and its class/quality of service capabilities.

a. Port Count

Server-based solutions are typically limited to 8-16 ports --a limitation of the server hardware platform. A switch can offer 12-24 at the low-end to hundreds at the high-end. This is important in applications where a large range of key resources needs to be protected; i.e., web server farm [Ref. 28].

b. Performance

Most firewalls are implemented based on an "us vs. them" scenario [Ref. 29]. In most cases, they sit between the WAN-connected external, untrusted world (Internet) and the LAN-oriented internal, trusted world. The

bottleneck is the WAN link--rarely over T1/E1. The need for access control is inside the network--on the LANs. Internal security devices must be able to keep up with LAN speeds--100Mb Ethernet. Widely deployed, wire-speed switches, already providing the network building blocks, are ideal platforms to implement these internal security control measures.

c. *Redundancy*

Rarely are firewalls implemented in redundant fashion. Most are dual-homed--based on internal-to-external paradigm [Ref. 29]. A switch can have hundreds of dedicated ports. Layer-three switches, which support Routing Information Protocol [Novell] (RIP) and Open Shortest Path First (OSPF), will provide route redundancy. Redundant interface cards and redundant chassis can be configured to provide a fail-safe backup scenario. Dual-power supplies will minimize other outage potential. Switches are built from the ground up with redundancy in mind.

d. *Management*

The few must manage the many. Therefore a more centralized policy administration is better [Ref. 29]. Security control measures that can be controlled from a central station simplify enterprise-wide security because policies are created centrally and distributed to wherever they are needed. Switches' running internal security control measures, though physically distributed, are centrally managed. With the new Simple Network Management Protocol (SNMP) v3 RFCs (2271-2275), secure network management based on authentication and encryption is not far away [Ref. 29].

e. QoS

All users and all traffic types are not equal in priority. Higher intelligence can segregate them. As distributed directory services begin to emerge, switches are ideally positioned to implement security policy once the user or application is identified. Smart businesses will see the value of putting first things first. The Oracle, SAP, or PeopleSoft databases are key business applications. These run the business--take the orders, move the inventory, collect the revenue and pay the expenses. Users getting a regular update on their stock portfolio should not be treated on the same level as a user running a business application [Ref. 23]. Incumbent switches are able to open these secure channels for the users and applications in real-time, at high speeds. See Table 2.

3. External Firewalls

External firewalls are positioned to control access in and out of a private network to a public network, namely the Internet [Ref. 30]. They regulate interactions between the trusted and untrusted networks. In a switch-based implementation, on one side, the firewall is connected via a WAN interface directly to the service provider (or sometimes a dedicated router is between the switch and the service provider). On the other side, the firewall is connected via a LAN interface to the private corporate network, often the backbone. Off to the side, for those hosting their own Internet applications, there will be a networking de-militarized zone (DMZ). The DMZ will typically contain one or more web servers, ftp servers, etc. The bandwidth requirement is mandated by the speed

of the Internet connection, typically one T1/E1 between networks.

Performance	20-30 Mbps NT; 70-80 Mbps Unix	25 Mbps today; 150 Mbps 1989, 1 Gbps 1999
Configuration	Minimal --redundant NIC's	Hardware & routes, paths
Management	Some --bandwidth management	Much --bandwidth management, RSVP, VLANs, applications, ATM
Hardware	Dedicated host system -- processing unit, memory/disk, monitor, keyboard, NICs	Integrated into existing platform which is providing other switch functionality
Applications	Internet (WAN-oriented)	Internet & Extranet (WAN-oriented); Intranet (LAN-oriented)

Table 2. A Switch-Based Firewall Has Two Primary Applications: as an External Firewall and as an Internal Firewall. [Ref. 30].

External firewalls are useful at securing main campus networks as well as remote offices or distributed facilities. See Figure 22.

External Firewalls

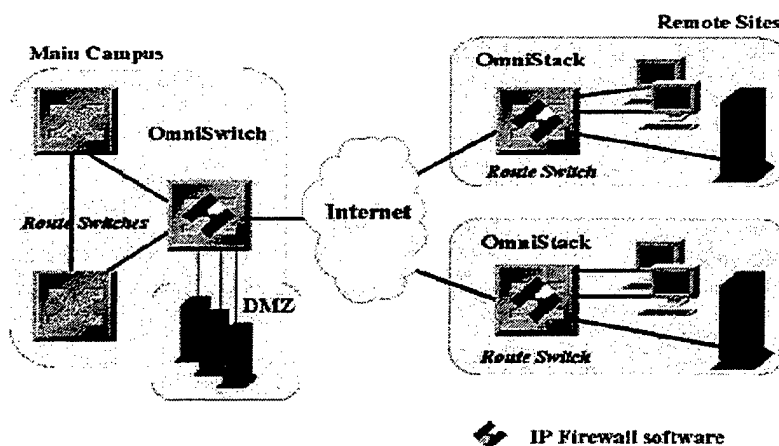


Figure 22. External Firewall. [Ref. 28]

4. Internal Firewalls

The firewall architecture in this section (Figure 23) is comprised of two primary modules: the Control Module and the Firewall Module, both explained later in this section. This firewall architecture example is used because it allows the definition of security policy in a central location and the distribution of it to all enforcement points. Multiple user access control allows different people across the organization to manage the security policy, based upon their authorization levels [Ref. 30]. Internal firewalls are deployed in locations inside organizations to protect key resources from internal threats. Internal firewalls are Intranet-oriented. Because statistics show that 60-80% of all security breeches resulting in financial loss are launched by insiders, protecting key resources from pseudo-trusted individuals is becoming critical [Ref. 30]. These insiders can be contract employees, disgruntled workers, or hackers that circumvented the Internet/external firewall and entered the network via a backdoor (modem, etc.). (Criminal hacking has a direct analogy with violent crime: the victims usually know their attackers.) Ideal locations for internal firewalls are between departments, between VLANs, at branch offices, in front of key physical resources, like a server farm [Ref. 30]. Internal firewalls are useful when isolating test labs, less secure/public corporate facilities, or "secret" subnetworks. Bandwidth requirements are much higher compared to external firewalls. Typically, they must be able to support Ethernet-to-Ethernet (10 or 100 Mbps, and shortly 1000 Mbps) or very high speed (eg. OG 48) ATM communications.

Internal Firewalls

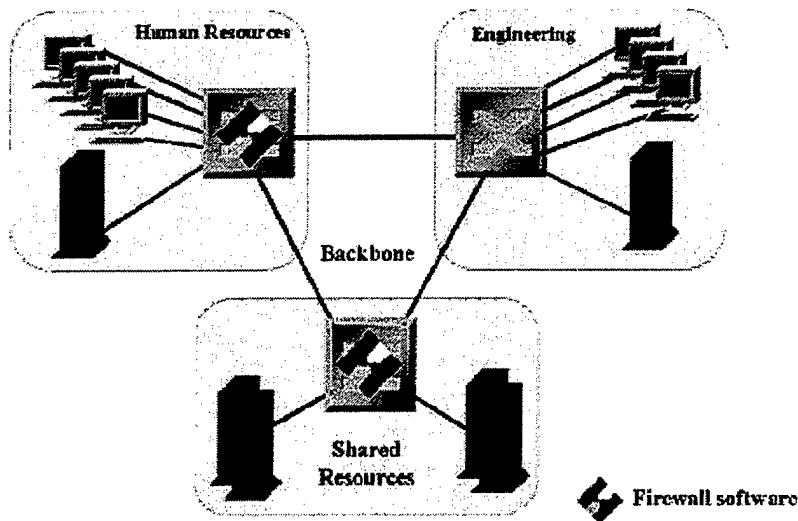


Figure 23. Internal Firewall. [Ref. 28]

5. Firewall Architecture

The **Control Module** includes the Graphical User Interface (GUI) and the Management Server. The GUI is the front-end to the Management Server, which manages the Firewall database--rule base, network objects, services, users, etc. The Control Module can be deployed in a client/server configuration. As an example of architectures, the client can run Windows 95, NT, or X/Motif GUI. It controls a Management Server running on Windows NT, SunOS, Solaris, HP-UX and AIX [Ref. 29].

The **Firewall Module** implements the security policy as defined by the Control Module. The Firewall Module communicates with the Control Module via daemons. See Figure 24. The components necessary to operate a switch-based firewall are [Ref. 29]:

- Management Server--known as the Enterprise Management Console (one Control Module can manage anywhere from 12-24 Inspection Modules, depending on the volume of traffic and the amount of logging required).
- GUI client (more than one GUI client can exist)
- Inspection Module--code that runs on each switch acting as a firewall.

6. Switch-embedded Firewall Features

A switch-embedded firewall packages the firewall in many different ways. Most of the packages are oriented towards standalone server-based implementations [Ref. 27]. An embedded firewall can be ported to run in a true networking device. (A true networking device is one that is built from the ground up for networking: optimized for large port counts, able to support multiple media types, protocols and services, and integrated redundancy. A

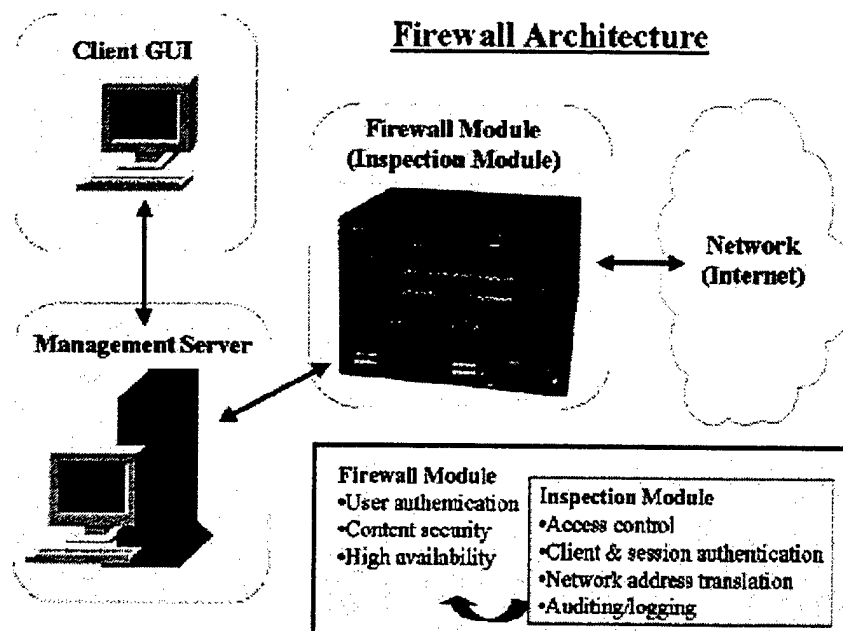


Figure 24. Firewall Architecture. [Ref. 27]

PC stuffed with interface cards would not be considered by most to be a true networking device.)

The package created for the embedded code is known as an Inspection Module. An Inspection Module is licensed per network device (switch). Each Inspection Module is capable of supporting four basic firewall functions [Ref. 29]:

- **Access control**
- **Logging**
- **Address translation**
- **Authentication**

a. Access control

Access control is the main role a firewall performs: enforcement of security policy for traffic that is routed through it. Policy rules are created to accept, reject, and log packet flows based on source, destination and service type. When implementing security policies on a firewall, the administrator must perform the following actions:

- Define the network objects to be used in the rule base
- Define any proprietary services used in the network
- Define the rule base for accepting, rejecting, and logging sessions
- Install the rule base on the gateways; i.e., switches

The network objects are the devices or networks to be secured. Policy rules are created based on the communication attempts: source, destination, service, and time (day, week). A service is the type of communication (TELNET, ftp, lotus notes, RealAudio, etc.). Check Point has

predefined over 160 different services. If a desired service is undefined, the firewall administrator must manually define it. Once the first four rule elements are bound, the system next must understand what to do with the packet(s): action (accept, drop, etc.) and track (whether or not to log the communication, generate an alert, etc.). Finally, the rule must state which network object; e.g., gateway/switch, will enforce the rule.

The firewall in Figure 23 follows the security principle of "all communications are denied unless expressly permitted [Ref. 29]." By default, this firewall drops traffic that is not explicitly allowed by the security policy and generates real-time security alerts, providing the system manager with complete network status. Each rule opens holes through the firewall for specific networks, devices, users, and services.

Common attacks are easily stopped by the Figure 23 firewall [Ref. 29]:

- IP Spoofing--an intruder attempts to gain unauthorized access by altering a packet's IP address to make it appear as though the packet originated in a part of the network with higher access privileges.
- Ping of Death--PING (ICMP) packets larger than 65508 are not handled well by kernels, making some systems crash or reboot.

Access to the firewall for management is distributed. Access levels are defined as follows:

- Read/Write: access to all functionality of the firewall's management tools.
- User Edit: modify user information only; access to all other functionality is read-only.

- Read Only: read-only access to the Security Policy Editor.
- Monitor Only: read-only access limited to the Log Viewer and the System Status tools.

(1) Authenticity--Firewall Authentication. Figure 23 firewall provides customers, including remote users and telecommuters, with secure, authenticated access to enterprise resources using multiple authentication schemes. The firewall's authentication securely validates that the users attempting to make a connection are who they say they are before the communication is allowed [Ref. 29]. Modifications to local servers or client applications are not required. All authentication sessions can be monitored and tracked through the Log Viewer. The demonstrated firewall provides three authentication methods: user, client, and transparent session.

(2) User Authentication. Provides access privileges on a per user basis for FTP, TELNET, HTTP, and RLOGIN, regardless of the user's IP address. Inspection Modules, such as Xylan's switch-based firewall implementation, do not support user authentication.

(3) Client Authentication. Enables an administrator to grant access privileges to a specific user at a specific IP address. In contrast to user authentication, client authentication is not restricted to specific services, but provides a mechanism for authenticating any application. It does not require any additional software or modifications on the client or server. The administrator determines how each user is authenticated, which servers/applications are accessible, what times/days, and how many sessions are permitted [Ref. 29].

(4) Transparent Session Authentication. It is used to authenticate any service on a per-session basis. After the user initiates a connection directly to the server, the firewall gateway, located between the user and the destination, intercepts the connection, recognizes that it requires user-level authentication, and initiates a connection with a Session Authentication Agent. The Agent performs the required authentication, after which the firewall allows the connection to continue to the requested server if permitted. This firewall supports the following authentication schemes [Ref. 30]:

- SecureID--The user is challenged to enter the number displayed on the SecureID card.
- S/Key--The user is challenged to enter the value of requested S/Key iteration.
- OS Password--The user is challenged to enter his or her OS password.
- Internal--The user is challenged to enter his or her internal firewall password on the gateway.
- Defender--The user is challenged for the response.
- RADIUS--The user is challenged for a response, as defined by the RADIUS server [31].

b. Logging

Figure 23 firewall allows the security manager to monitor accounting data on selected connections. For each connection handled by the rule, an accounting log entry is generated which includes a dozen or so fields including the connection's duration, the number of bytes and the number of packets transferred. Log records are generated when the monitored connection

ends, so they can be viewed in the Log Viewer. The Active Connections View can be used to monitor ongoing connections. The live connections are stored and handled in the same way as ordinary log records, but are kept in a special file that is continuously updated as connections start and end. All the standard Log Viewer features, such as selection, search engine, etc., can be used to monitor current network activity. When using the accounting option, the connection accounting data (time elapsed, bytes and packets transferred) is continuously updated, so the security manager can monitor not only the fact of the connection but also its activity [Ref. 30].

The firewall provides integration of multiple alert options including email notification and SNMP traps for integration with SNMP-based network management systems such as HP OpenView, SunNet Manager, or IBM's NetView 6000 [Ref. 30]. A User Defined alerting mechanism is available to integrate with paging, trouble ticketing and help desk systems providing additional flexibility [Ref. 32].

c. Address Translation

The firewall's Network Address Translation feature conceals internal network addresses from the Internet, avoiding their disclosure as public information. This feature overcomes IP addressing limitations, including restricted IP address allocation and unregistered internal addressing schemes [Ref. 32]. The firewall maintains the integrity of an organization's internal addressing scheme by mapping unregistered IP addresses with valid ones for full

Internet access. There are two modes of operation -- dynamic mode and static mode.

(1) Dynamic Mode. Provides users access to the Internet while conserving registered IP addresses and hiding the actual IP addresses of network resources. Dynamic mode uses a single IP address to map all connections through the protected access point. Since the IP address used in dynamic mode is used only for outbound communication and not by any resource, there is nothing to hack or spoof. The firewall allows an unlimited number of addresses to be dynamically mapped to a single IP address [Ref. 32].

(2) Static Mode. Provides a one-to-one assignment between the published IP address and the real IP address. Static mode would typically be implemented when administrators did not wish to expose the real IP addresses of the network servers, or if a network IP address had been assigned historically and they needed to provide "real" addresses so that people on the Internet can access them [Ref. 32].

There are two methods for specifying address translation [Ref. 32]. The first is to specify automated address translation during the object definition process. This will automatically generate the appropriate translation rule. The second is to specify the address translation specifications using the address translation rules editor. All network objects can be used to specify address translation rules. The Figure 23 firewall has the ability to validate the specified address translation rules, helping to avoid configuration mistakes [Ref. 33].

G. LOOKING AHEAD--DIRECTORY INTEGRATION

Organizations will begin to move to an enterprise-wide directory structure for providing policy-based control of network resources. Through secure authentication techniques, the network will begin to provide users and applications a certain quality of service. Users will sign on once, regardless of their location, and be able to access all of their authorized network resources. The network infrastructure--the intelligent multi-layer switches--will integrate directory-enabled applications. The distributed architecture will be able to dynamically modify the available network resources based on a user's identity [Ref. 33].

H. TYPES OF COMPUTER SECURITY POLICY

This section will discuss four types of computer security policy, their components, and aspects of policy implementation. Program-level policy is used to create an organization's computer security program. Program-framework policy establishes the organization's overall approach to computer security (i.e., its computer security framework). Issue-specific policies address specific issues of concern to the organization. Lastly, system-specific policies focus on policy issues which management has decided for a specific system. Comparison of an organization's computer security policies to the types described in this bulletin will assist managers in determining if their policies are comprehensive and appropriate [Ref. 26].

1. Program-level Policies

Organizations need program-level policy to establish the security program,

assign program management responsibilities, state organization-wide computer security purpose and objectives, and provide a basis for compliance. Program-level policy is typically issued by the head of the organization or another senior official, such as the top management officer [Ref. 26].

2. Program-framework Policies

These policies provide organization-wide direction on broad areas of program implementation. For example, they may be issued to assure that all components of an organization address contingency planning or risk analysis. They are appropriate when an organization can yield benefits from a consistent approach. Program-framework policies are issued by a manager with sufficient authority to direct all organization components on computer security issues. This may be the organizations management official or the head of the computer security program [Ref. 26].

3. Issue-specific Policies

These identify and define specific areas of concern and state the organizations position. Depending upon the issue and attendant controversy, as well as potential impact, issue-specific policy may come from the head of the organization, the top management official, the Chief Information Officer, or the computer security program manager [Ref. 26].

4. System-specific Policies

State the security objectives of a specific system, define how the system should be operated to achieve the security objectives, and specify how the protections and features of the technology will be used to support or enforce the

security objectives. A system refers to the entire collection of processes, both automated and manual. System-specific policy is normally issued by the manager or owner of the system (which could be a network or application), but may originate from a high official, particularly if all impacted organizational elements do not agree with the new policy [Ref. 26].

I. EXAMPLES OF SECURITY POLICY

1. Program-level Policy

Program-level policy establishes the computer security program and its basic framework. This high-level policy defines the purpose of the program and its scope within the organization, assigns responsibilities for direct program implementation (to the computer security organization) as well as responsibilities to related offices (such as the IRM organization), and addresses compliance issues. Components of program-level policy should include [Ref. 26]:

- **Purpose:** Clearly states the purpose of the program. This includes defining the goals of the computer security program as well as its management structure. Security-related needs, such as integrity, availability, and confidentiality, can form the basis of organizational goals established in policy. For instance, in an organization responsible for maintaining large mission-critical databases, reduction in errors, data loss, or data corruption might be specifically stressed. In an organization responsible for maintaining confidential personal data, however, goals might emphasize stronger protection against unauthorized disclosure. The program management structure should be organized to best address the goals of the program and respond to the particular operating and risk environment of the organization. Important issues for the structure of the central computer security program include management and coordination of security-related resources, interaction with diverse communities, and the ability to relay issues of concern to upper management. The policy could also establish operational security offices for major systems, particularly those at high risk or most critical to organizational operations.

- **Scope:** Specifies which resources (including facilities, hardware, and software), information, and personnel the program covers. In many cases, the program will cover all systems and agency personnel, but this is not always true. In some instances, a policy may name specific assets, such as major sites and large systems. Often tough management decisions arise when defining the scope of a program, such as determining the extent to which the program applies to contractors and outside organizations utilizing or connected to the organization's systems. The Computer Security Act of 1987 requires federal agencies to address the security of all federal interest systems.
- **Responsibilities:** Addresses the responsibilities of officials and offices throughout the organization, including the role of line managers, applications owners, users, and the data processing or IRM organization. The policy statement should distinguish between the responsibilities of computer services providers and the managers of applications utilizing the computer services. It can also serve as the basis for establishing employee accountability. Overall, the program-level assignment of responsibilities should cover those activities and personnel who will be integral to the implementation and continuity of the computer security policy.
- **Compliance:** Authorizes the use of specified penalties and disciplinary actions for individuals who fail to comply with the organization's computer security policies. Since the security policy is a high-level document, penalties for various infractions are normally not detailed here. However, the policy may authorize the creation of compliance structures which include violations and specific penalties. Infractions and associated penalties are usually defined in issue-specific and system-specific policies. When establishing compliance structures, consider that violations of policy can be unintentional on the part of employees. For example, nonconformance can be due to a lack of knowledge or training.

a. *An Example of a Program Level Policy*

The information residing on the XYZ Agency local area network (LAN) is mission critical. The size and complexity of the LAN within XYZ has increased and now processes sensitive information. Because of this specific security measures and procedures must be implemented to protect the

information being processed on the XYZ LAN. The XYZ LAN facilitates sharing of information and programs by multiple users. This environment increases security risk and requires more stringent protection mechanisms than would be needed for a standalone microcomputer (PC) operation. These expanding security requirements in the XYZ computing environment are recognized by this policy which addresses the use of the XYZ LAN.

This policy statement has two purposes. This first is to emphasize for all XYZ employees the importance of security in the XYZ LAN environment and their role in maintaining that security. The second is to assign specific responsibilities for the provision of data and information security, and for the security of the XYZ LAN itself.

2. Program-Framework Policy

Program-framework policy defines the organization's security program elements which form the framework for the computer security program and reflect decisions about priorities for protection, resource allocation, and assignment of responsibilities [Ref. 26]. Criteria for the types of areas to be addressed as computer security program elements include, but are not limited to [Ref. 26]:

- Areas for which there is an advantage to the organization by having the issue addressed in a common manner;

- Areas which need to be addressed for the entire organization; areas for which organization-wide oversight is necessary;

- Areas which, through organization-wide implementation, can yield significant economies of scale.

The types of areas addressed by program-framework policy vary within each organization as does the way in which the policy is expressed. Some organizations issue policy directives, while others issue handbooks which combine policy, regulations, standards, and guidance. Many organizations issue policy on key areas of computer security, such as life cycle management, contingency planning, and network security.

Keep in mind the criteria stated above for the types of areas that should be addressed in program-framework policy. If the policy (and its implementing standards and guidance) is too rigid, cost-effective implementations and innovation could be stifled.

a. Program-Framework Policy Example

As an example of program-framework policy, consider a typical organization policy on contingency planning. The organization might require that all contingency plans categorize criticality of processing according to a standard scale. This will assist the organization in the preparation of a master plan (for use if the organization's physical plant is destroyed) by facilitating prioritization across intra-organizational boundaries. Policy in these areas normally applies throughout the organization and is usually independent of technology and the system or application. Program-framework policies may be comprised of components similar to those contained in program-level policy—but may be in a very different format (e.g., in organizational handbook directives).

3. Issue-Specific Policy

Issue-specific policies focus on areas of current relevance and concern (and sometimes controversy). Program-level policy is usually broad enough that it requires little modification over time. Conversely, issue-specific policies require more frequent revision due to changes in technology and related factors. As new technologies develop, some issues diminish in importance while new ones continually appear [Ref. 26].

It may be appropriate, for example, to issue a policy on the proper use of a cutting-edge technology or problem, the security vulnerabilities of which are still largely unknown. A useful structure for issue-specific policy is to break the policy into its basic components: statement of an issue, statement of the organization's position, applicability, roles and responsibilities, compliance, and points of contact. Other topic areas may be added as needed.

- **Issue Statement:** Defines the issue, with any relevant terms, distinctions, and conditions. For example, an organization might want to develop an issue-specific policy on the use of unapproved software, which might be defined to mean any software not approved, purchased, screened, managed, and owned by the organization. Additionally, applicable distinctions and conditions might need to be included, for instance, software privately owned by employees but approved for use at work and for software owned and used by other businesses under contract to the organization.
- **Statement of the Organization's Position:** Clearly states the organization's position on the issue. To continue the example of unapproved software, the policy would state whether use of unapproved software is prohibited in all or some cases, whether or not there are further guidelines for approval and use, or whether case-by-case exceptions will be granted, by whom, and on what basis.

- **Applicability:** Clearly states where, how, when, to whom, and to what a particular policy applies. For example, the hypothetical policy on unapproved software may apply only to the organization's own on-site resources and employees and not to contractor organizations with offices at other locations. Additionally, the policy's applicability to employees travelling among different sites or working at home who will transport and use disks at multiple sites might require clarification.
- **Roles and Responsibilities:** Assigns roles and responsibilities. To continue the software example, if the policy permits unapproved software privately owned by employees to be used at work with appropriate approvals, then the approving authority would be identified. An office responsible for compliance could also be named.
- **Compliance:** Gives descriptions of the infractions which are unacceptable and states the corresponding penalties. Penalties must be consistent with organizational personnel policies and practices and need to be coordinated with appropriate officials, offices and, perhaps, employee bargaining units.
- **Points of Contact and Supplementary Information:** Gives the name of the appropriate individuals to contact for further information and lists any applicable standards or guidelines. For some issues the point of contact might be a line manager; for other issues it might be a facility manager, technical support person, or system administrator. For yet other issues, the point-of-contact might be a security program representative. Using the software example, employees need to know whether the point of contact for questions and procedural information would be the immediate superior, a system administrator, or a computer security official.

a. *Issue-Specific Policy Example*

All new Products purchases, including software purchased to adapt present applications to Year 2000 compliance, and all other Products installed in this information system must comply with the provisions of the following "Year 2000 Warranty" [Ref. 23]:

- Seller or Licensor (which term shall also include all persons designing, developing or installing software or other Products with

or without an express agreement with the organization) represents and warrants that the Products are designed to be used prior to, during, and after the calendar year 2000 A.D., and that the Products will operate during each such time period without error relating to date data, specifically including any error relating to, or the product of, date data which represents or references different centuries or more than one century

- Without limiting the generality of the foregoing, the Seller or Licensor further represents and warrants: That the Products will not abnormally end or provide invalid or incorrect results as a result of date data, specifically including date data which represents or references different centuries or more than one century
- That the Products have been designed to ensure year 2000 compatibility, including, but not limited to, date data century recognition, calculations which accommodate same century and multi-century formulas and date values, and date data interface value that reflect the century
- That the Products include "year 2000 capabilities" which means they:
 - will manage and manipulate data involving dates, including single century formulas and multi-century formulas, and will not cause an abnormally ending scenario within the application or generate incorrect values or invalid results involving such dates.
 - will provide that all date-related user interface functionalities and data fields include the indication of century
 - will provide that all date-related data interface functionalities include the indication of century
 - will provide for accurate processing of date/time data (including, but not limited to, calculating, comparing, and sequencing) from, into, and between centuries, i.e. the twentieth and twenty-first centuries.
 - when used in combination with other information technology, shall accurately process date/time data if the other information technology properly exchanges date/time data with the Products.

This organization shall not enter into any Products agreement with provisions which tend to limit or eliminate the liability of any party with respect to the Year 2000 Warranty above. Remedies providing for reasonable liquidated damages for failure of the Year 2000 Warranty shall be acceptable. Any purchase or other use of Products not in compliance with this policy shall not be permitted without the express consent of the administrative unit's director of information systems.

4. System-Specific Policy

Program-level policy and issue-specific policy both address policy from a broad level, usually encompassing the entire organization. System-specific policy, on the other hand, is much more focused, since it addresses only one system. Many security policy decisions apply only at the system level [Ref. 26].

Some examples include:

- User X may or may not be allowed to read or modify data in the system.
- User X may or may not depending on the conditions read or modify data.
- User X may or may not be allowed to dial into the computer system from home or while on travel.

To develop a comprehensive set of system security policies, use a management process which derives security rules from security goals. A hypothetical company, ABC Sales, has a headquarters office in San Francisco and 100 branch offices around the U.S. ABC also has traveling sales representatives who use portable computers. Employees within the company

have Internet access for Web browsing, e-mail, and other uses. There's also a public Web site, hosted at the main office. Vendors and customers exchange e-mail and sometimes even large files with the company.

For ABC, mobility is one of the big requirements. If the CEO travels to the Los Angeles sales office, he or she might want to print a document using a laptop. This means giving the CEO temporary printer access on the Los Angeles network. Typically, there might be a guest cubicle in the sales office, preconfigured with a guest account and network connection. But this isn't the best setup since it allows network access (however limited) to anyone who can reach the guest cubicle. So the CEO might instead use a smart card that ensures the same access rules apply wherever they plug in. Of course, the use of smart cards requires a unified policy database—and devices that implement the unified policy.

Meanwhile, back in San Francisco, employees should be blocked from using the CEO's desktop PC to surf the Web. Centralized policy can help here as well. For example, the network might employ IPsec's user certificates to control access to an external network proxy server. In this scenario, a policy could grant multiple levels of access, like giving full access to executives and more limited, job-related access to other employees. Consider a three-level model for system security policy [Ref. 26]:

- Security Objectives
- Operational Security
- Policy Implementation.

a. *Security Objectives*

First, define security objectives. While this process may start with an analysis of the need for integrity, availability, and confidentiality, it cannot stop there. A security objective must be more specific, concrete, and well-defined. It also should be stated so that it is clear that the objective is achievable. The security objectives should consist of a series of statements which describe meaningful actions about specific resources. These objectives should be based on system functional or mission requirements, but should state the security actions which support the requirements.

b. *Operational Security*

Next lay out the operational policy which gives the rules for operating a system. Following the same integrity example, the operational policy would define authorized and unauthorized modification: who, (by job category, by organization placement, or by name) can do what (modify, delete, etc.) to which pieces of data (specific fields or records) and under what conditions. Managers need to make decisions in developing this policy since it is unlikely that all security objectives will be fully met. Cost, operational, technical, and other constraints will intervene. Consider the degree of granularity needed for operational security policies. Granularity refers to how specific the policy is with regard to resources or rules. The more granular the policies, the easier to enforce and to detect violations. A policy violation may indicate a security problem. In addition, the more granular the policy, the easier to automate policy enforcement. Consider the degree of formality you want in documenting the

policy. Once again, the more formal the documentation, the easier to enforce and to follow policy. Formal policy is published as a distinct policy document; less formal policy may be written in memos. Informal policy may not be written at all. Unwritten policy is extremely difficult to follow or enforce. On the other hand, very granular and formal policy at the system level can also be an administrative burden. In general, good practice suggests a granular formal statement of the access privileges for a system due to its complexity and importance.

Documenting access controls policy makes it substantially easier to follow and to enforce. Another area that normally requires a granular and formal statement is the assignment of security responsibilities. Some less formal policy decisions may be recorded in other types of computer security documents such as risk analyses, accreditation statements, or procedural manuals. However, any controversial, atypical, or uncommon policies may need formal policy statements. Atypical policies would include any areas where the system policy is different from organization policy or from normal practice within the organization, either more or less stringent. They should also contain a statement explaining the reason for deviation from the organization's standard policy.

c. Policy Implementation

Determine the role technology will play in enforcing or supporting the policy. Security is normally enforced through a combination of technical and traditional management methods. While technical means are likely to include the use of access control technology, there are other automated means of enforcing or supporting security policy. For example, technology can be used to block

telephone systems users from calling certain numbers. Intrusion detection software can alert system administrators to suspicious activity or take action to stop the activity. Personal computers can be configured to prevent booting from a floppy disk.

Automated security enforcement has advantages and disadvantages. A computer system, properly designed, programmed, and installed, consistently enforces policy, although no computer can force users to follow all procedures. In addition, deviations from the policy may sometimes be necessary and appropriate. This situation occurs frequently if the security policy is too rigid.

J. THE SECURITY CHALLENGE

Formulating viable computer security policies is a challenge for an organization and requires communication and understanding of the organizational goals and potential benefits to be derived from policies. Through a carefully structured approach to policy development, which includes the delegation of program management responsibility and an understanding of program-level, program-framework, issue-specific, and system-specific policy components, your organization can achieve a coherent set of policies. These will help produce a framework for a successful computer security program [Ref. 26].

THIS PAGE INTENTIONALLY LEFT BLANK

V. NETWORK COST MANAGEMENT POLICY: (A CASE STUDY)

A. EXECUTIVE SUMMARY

There is an abundance of new vendor solutions vying for today's networking business [Ref. 16]. Typically those who oversee their organization's network are skeptical at best. It is not that they are unimpressed with ATM, Layer 3 switching, IP Multicast, Gigabit Ethernet and other next-generation solutions; it is more that they are too busy trying to justify the hefty price tag to keep the present network up and running [Ref. 34]. And while many new products bring desirable features like increased bandwidth, per-port RMON and broadcast control to evolving infrastructures, do they really remedy the problems associated with the bottom-line success?

After researching the subject of cost policy, it is the the author's opinion that the real concern heading into the future is not necessarily migrating to high-speed intranets and extranets, but reducing the spiraling costs associated with supporting and operating the current network [Ref. 34].

It may be concluded that the network is a business entity that must be preserved. There are strategies that can be implemented to reduce the complexity of the enterprise and lower overall costs [Ref. 34]. The only problem facing senior managers is choosing which path to follow to reap such rewards [Ref. 34]. In this chapter, the author will address some of the major factors that drive network costs up, and then provide a working example of how to best

overcome these factors while still building a productive, profitable IT infrastructure.

In developing this chapter, it was found that there are two distinct strategies available to organizations: a primary vendor strategy that essentially relies on a single source for all networking components, and a multi-vendor strategy that draws from the expertise of several leaders in their respective fields. Both schools of thought have their merits, but it is believed here that by adopting a best-of-breed approach where one can leverage the strengths of different vendors, one can better reduce network costs as the foundation is laid for a truly integrated enterprise that will accommodate all converging resources.

B. NETWORK COST PROBLEM DEFINED

When addressing the costs of running a corporate network, most analysts or research groups tend to break down the costs into three distinct categories: capital equipment, support staff and facilities [Ref. 34]. Strategic Networks Consulting (SNC), in its paper "Operating Today's Corporate Network," reports that in an average three-year span, a company's IT budget will devote 29% for capital equipment, 37% for staff and 34% for facilities. Facilities are generally defined as those services that maintain the physical and logical well-being of the network. For instance a facility cost associated with the infrastructure could include leased line services or operational health services such as hardware, software and circuit maintenance.

Most organizations tend to spend more on the day-to-day maintenance of the network than capital equipment or support staff [Ref. 34]. But when looking

even closer at these categories, one finds that there may be some other variables that account for the associated costs.

C. BANDWIDTH ISSUES

The outlay for network equipment is substantial but at least it is quantifiable. That is, it used to be, until bandwidth demands dictated replacing devices with faster, smarter solutions [Ref. 34].

Network bandwidth is a precious commodity for the individual user. It is what enables the user to access various applications, download selected images, and exchange information with their neighbor next door or across the world. But as the applications have evolved and grown much larger and much more impressive, the need for bigger pipes to deliver them has also grown [Ref. 1].

Three years ago, shared 10 Mbps to the desktop was more than adequate for the average user. Now with the prevalence of multimedia networked applications, streaming video, 3D CAD/CAM programs and such, end users are requiring a significant increase in their allotted bandwidth. It is not uncommon to see some engineering departments relying on a dedicated 100 Mbps or higher Full Duplex connection to the desktop to run these heavy applications [Ref. 1].

Even your typical Web surfers who need to go out on the Internet for valuable company research need faster transmission to download images and clips. You certainly can not deny the work-force the freedom and flexibility to run these applications and browsers; but many organizations are asking when the demand is going to stop [Ref. 34]. Does it mean having to replace equipment every one or two years? And what about consolidating information resources

such as voice, video and data? Managers must ask is their present infrastructure capable of such integration. It would be nice to know that once invested in a solution, it would stick around for at least the next five years. Unfortunately, this expense never seems to stop coming back to haunt organizations [Ref. 34].

D. CONFIGURATION ISSUES

In a survey conducted by SNC, seven enterprise customers were asked which tasks consumed most of their time and, not surprisingly, the number one operational task mentioned was configuration [Ref. 35]. SNC defined configuration as the re-configuring of devices as a result of user moves, adds and changes including setting router filters to govern the flow of user traffic or changing the way a router handles a particular protocol.

The report went on to say that "because configuration accounts for so much support staff time, and people costs incurred during the operational life cycle phase account for the greatest percentage of overall cost of network ownership, configuration tasks are extremely expensive...specifically, configuration tasks account for eight to nine percent of the typical users' total cost of network ownership over a three year period [Ref. 35]."

In dollar amounts, Strategic Networks estimates that configuration can cost a company an average of \$250,000 to \$300,000 per year based on an annual network budget of \$2.9 million. These figures were then broken down to a per-move basis and it was determined that each move cost between \$250 and \$400. If companies are going to remain profitable, they certainly can't afford to stand still. Unfortunately, there is a price to be paid for all this moving around.

E. DOWNTIME DILEMMA

Interestingly, the second most mentioned operational task mentioned in SNC's survey was fault management, including resolving problems that arise from network equipment malfunctions. This is usually performed in direct response to user complaints or network management system alarms. As covered in Chapter III while earlier network implementations were relatively easy to troubleshoot, today's more complex infrastructures make pinpointing errors a challenging and, more importantly from a cost standpoint, a time consuming task.

This presents a two-fold cost problem. First, the MIS staff is too busy putting out fires instead of dealing with plans and strategies that would create a smoother running infrastructure. If they are spending all this time performing mundane tasks, does this require the hiring of more people with high salaries to help out? If not, how can one expect to migrate the network to a more cost-effective enterprise when there is no one available to lead the charge [Ref. 34]?

What also must be considered is what downtime does to the network users and customers. Now more than ever, companies simply cannot afford even the slightest disruption to the network. Some estimates put lost revenue in the millions of dollars for every minute of network downtime. This of course depends on the exact outage and the number of users affected, but suffice it to say, regardless of how large or small your organization is, any delay in service is one too many [Ref. 35].

F. SECURITY ISSUES

While not always considered a direct cost in terms of dollars and cents, network security is still an important issue when it comes to cost of ownership [Ref. 23]. After all, any breach of network security can have a profound effect on the future of an organization. In addition, the lack of policies within a corporate network opens the door for abuse of bandwidth. Maybe the horror stories of departments playing Doom over the network for endless hours on company time bring the story home even more [Ref. 23].

But while security is necessary, it too does not come cheaply. Most of the cost is buried in the configuration and integration of security features into the network [Ref. 34]. The capital costs are minimal compared to the manpower needed to design and implement an enterprise-wide network security program. We have already covered how much time and money is spent when stretching an MIS staff too thin.

A recent study conducted by the Registry, Inc. was able to put a dollar amount on security stating that mid-range companies with disparate, non-integrated network security systems spend nearly \$250 per year per networked desktop in support costs [Ref. 34]. Such costs should be addressed if a company is really committed to reducing the cost of running a tight, all-encompassing enterprise [Ref. 36].

G. AN ORGANIZATIONAL PLAN OF ACTION

Without the help of a crystal ball, organizations are hoping that their existing infrastructure will be able to support all the applications and services

they will require in the years ahead. If not, it is safe to say that capital expenses will give the day-to-day operational costs a run for the money in terms of which costs a company more [Ref. 34].

Already organizations are anticipating the need to consolidate voice, video and data on their enterprise. But how easily this integration will be is still tough to judge since vendors are still posturing on the subject without a clear, proven winner leading the charge. Needless to say, however, most companies are going to have to brace for some upgrade. How much depends on their vendor or vendors of choice.

The spiraling costs associated with running an enterprise network apply to everyone, even those companies that are in the business of building networks. Cabletron Systems, for instance, was experiencing many of the same growing pains on its own corporate net [Ref. 34]. There was not enough bandwidth to go around. Support staff was too busy performing day-to-day configuration tasks. Network delays were costing departments critical access to resources.

Fortunately, the company did not have to look far for a viable solution. By going to a full scale switched infrastructure (where routing functions would be used only when necessary) and implementing a unique set of performance-enhancing applications to complement a distributed management platform, Cabletron has been able to dramatically reduce costs as it increases reliability and productivity company-wide. The company's approach to the problem can be applied to many customers' infrastructures. Starting at the single-site migration and then moving to a multi-site, cross-country solution:

1. A Single-Site Solution

Since 1994, Cabletron has implemented a switch-based network solution at several of its key offices, but if the evolution of its corporate headquarters in Rochester, NH is targeted, we can better gauge the results of migrating to this type of solution, see Figure 25, [Ref. 34]. This case study may also be more applicable to customers since a good percentage of existing LAN environments are similar to Cabletron's infrastructure in terms of topologies, number and types of users, and capital equipment.

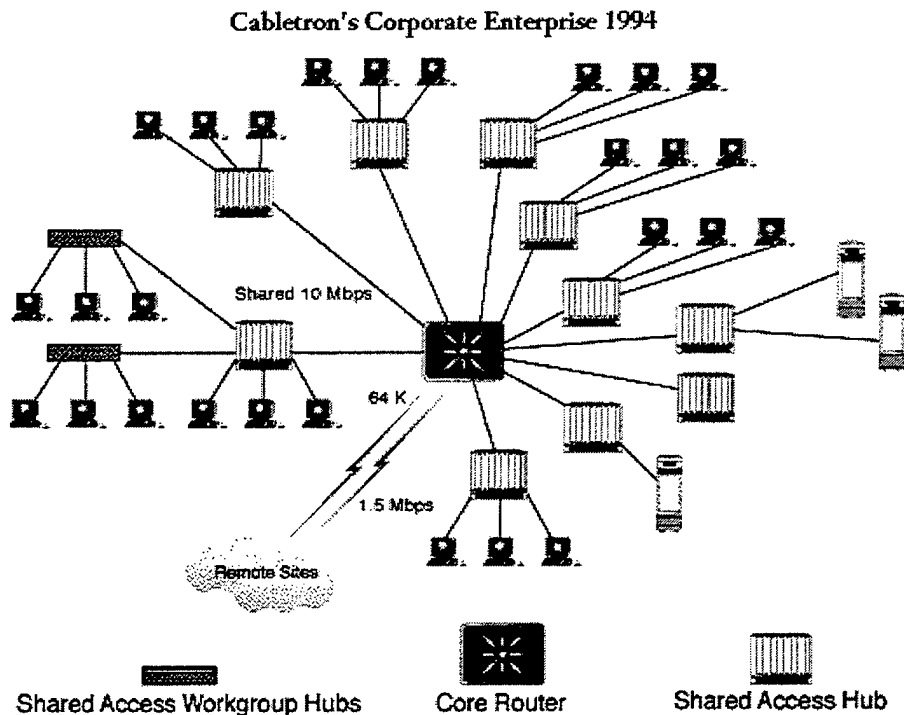


Figure 25. Corporate Enterprise 1994. [Ref. 34]

Cabletron's earlier network, with a router at the core, used a collapsed routed backbone topology to support two FDDI rings, 24 Ethernet subnets, and eight serial interfaces. This rather standard configuration served Cabletron well

in the days of shared access. But as user demands grew, Cabletron's Network Management group began integrating switching products into the existing network [Ref. 34].

In September 1994, the first MMAC-Plus SmartSwitch chassis (now called the SmartSwitch 9000), with an FDDI switching module, was installed in the main building's wiring closet, delivering 100 Mbps of unimpeded performance to the backbone. But that was just the beginning.

Through January 1995, Cabletron's Network Management group gradually integrated more and more switching products into the building's existing infrastructure including several of the wiring closets. Based around a switched FDDI backbone, the new MMAC-Plus switched network implemented all new products, with the exception of a traditional shared access hub in each wiring closet to support legacy equipment.

Tom Dunigan, Cabletron's corporate network manager, explained that as the switched network continued to grow, the core router became saddled more and more with heavy network traffic.

As the company's sub-networks expanded, particularly those groups within Engineering, the router couldn't handle the busy flow of traffic. More and more engineers were using their workstations as file servers, thereby putting stress on the router to filter and forward large amounts of data. Consequently, the router's performance decreased, as broadcasts and bottlenecks became commonplace [Ref. 34].

In other words, some of the main factors that helped drive up the cost of owning the network (i.e. inadequate throughput and network downtime) still existed. The idea of removing the router was a radical idea at the time due to the

router's essential but complicated role. Before you could even think about such a measure, several issues had to be resolved. Cabletron's Network Management group had to research the many side-effects that would inevitably emerge upon displacing the core router. One such issue was how to handle subnet masking.

As John Smart, senior engineer, explained,

There are two major protocols that our network uses -- TCP/IP and Novell's IPX. Within these protocols, you have subnet numbers, subnet masks, default gateways and network numbers to deal with. Before we could replace our router with switched-based products, we needed to aggregate various subnets -- which up to this point had relied on the router to filter and forward data -- and place those subnets onto switching boards within the MMAC-Plus switches [Ref. 34].

To accomplish this, however, every workstation must use the same number of bytes in their IP subnet masks. All Cabletron workstations use two-byte subnet masks, so placing those workstations on switching boards didn't present a problem to the Network Management group.

Dunigan said the most important element to successfully removing routers from the centers of networks is to have a strong knowledge of what the router can and can't do in a network.

There is tremendous physical and administrative work required with routed networks. Our department spent most of its time troubleshooting protocol-related problems and subnet issues [Ref. 34].

Smart agreed.

Highly-skilled technicians are needed to keep the routers going. You can easily collapse large segments of a network with one router-related mistake, affecting the entire operation of a company [Ref. 34].

As the migration continued, the network's reliance on the router steadily decreased. The MMAC-Plus hubs, with their switching management boards, assumed total responsibility for filtering and forwarding data. In January 1995, the router was completely removed from the network. According to Dave Dickson, network systems engineer,

We shut the router down for good. It was a gigantic step for us, and the company as a whole. We were making the final leap to a completely switch-based network [Ref. 34].

Running 200 Mbps of bandwidth, high-end Auspex file servers transfer network data across FDDI rings. Although strictly an FDDI-switched network, there is one ATM link between Cabletron's data center and an MMAC-Plus switch in a nearby facility. Fourteen FDDI interfaces run a total of 1.4 Gbps of data throughput across the MMAC-Plus' backplanes. Routers still remain at the outer edges of the network to connect Cabletron to the Internet and other remote networks. This solution provides more than enough bandwidth to selected departments and users.

High throughput aside, Cabletron's switched network has something else working in its favor. Every SmartSwitch supports the company's innovative SecureFast, a standards-based strategic architecture for building private and public networking infrastructures [Ref. 34]. Drawing from Cabletron's expertise in switching and enterprise management, SecureFast not only enables SmartSwitches to bypass the core router, but provides advanced tools and functionality that help create a cohesive, utility-like network that is consistent with the long-term goals of Cabletron's business [Ref. 34].

With the ability to deliver critical end-user information to those in charge of specific departments within Cabletron, SecureFast helps strip away many of the technical, complicated aspects associated with the network. Sure, the technology behind the hubs, switches and management system remains complex, but with SecureFast implemented, no one within Cabletron except MIS or network administrators needs to think about the bits and bytes (although SecureFast can significantly ease their jobs as well). Non-technical people simply have confidence the system works and can access these business-oriented applications right from their desktop [Ref. 34].

SecureFast is tightly integrated with Cabletron's SPECTRUM enterprise-wide management platform, and together they oversee the performance of the entire network. SPECTRUM's industry-leading, client/server architecture allows for the distribution of management tasks throughout the network, while also providing redundancy and minimizing expensive WAN bandwidth usage. This three-pronged solution, including SmartSwitches, SecureFast and SPECTRUM, has led to some immediate and long-term improvements in cost of ownership [Ref. 34].

Since the implementation of the new infrastructure and similar network upgrades at other sites, Cabletron's Network Management department has received positive feedback from employees and department heads alike. In terms of cost of ownership, the network has realized these benefits [Ref. 34]: Increased bandwidth, increased productivity. "Bandwidth utilization has improved dramatically," John Smart said. "And we now get unsolicited phone

calls from many of our engineers and managers about how fast the network is. "As Cabletron continues to grow, the network will no longer feel the strain of adding more users." "Bottlenecks have been completely eliminated at the center of our switch-based network," Dunigan said [Ref. 34].

Smart adds,

Larger impact on our available bandwidth was inevitable. Engineers were using more powerful workstations to compile their source code, CAD developers were eating up bandwidth with their mechanical and component designs, and even word processing jobs were taxing the network due to the large number of new users hooking into our network on a weekly basis [Ref. 34].

Cabletron's upgraded switched network now provides an average of 10 Mbps per user with some departments and workgroups getting up to 100 Mbps Full Duplex piped to the desktop. This increase not only eliminates any delays in using bandwidth-heavy programs like 3D CAD/CAM drawing but it enables these engineers to crank out work at a much faster pace, greatly improving productivity department-wide, see Figure 26.

As far as future bandwidth requirements are concerned, Cabletron's network has that covered too, see Figure 27. The distributed switching architecture of the SmartSwitches enables the company to add bandwidth as it adds modules. This not only protects the long-term investment in the switching chassis, but gives the company the flexibility to mix and match technologies as its growing needs dictate. For instance, if Cabletron wants to increase wide area access for certain departments within the corporate net, it simply can add a frame relay switch module to the SmartSwitch 9000. Or if the engineering group decides it requires even bigger pipes to support essential applications, a Gigabit

Ethernet switch can be installed. This kind of flexibility keeps capital investment costs down, even when integrating voice and video.

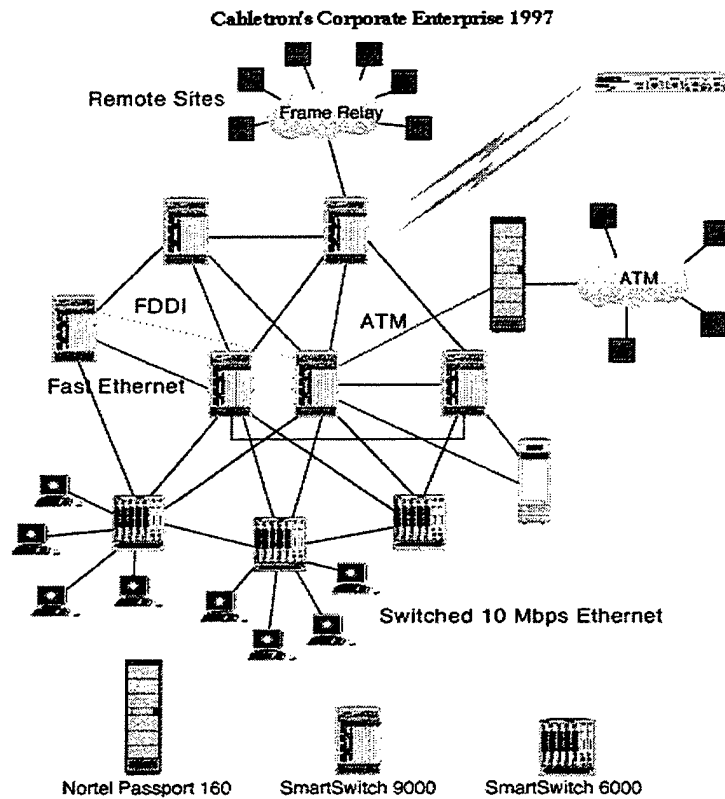


Figure 26. Corporate Enterprise 1997. [Ref. 34]

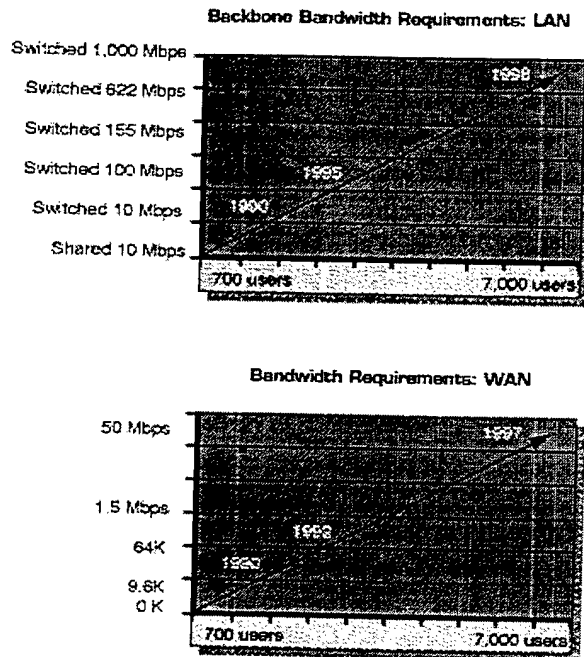


Figure 27. Bandwidth Graphs. [Ref. 34]

As emphasized in the Strategic Networks Consulting study, configuring the network is a number one concern for network administrators, accounting for almost \$350 per move when you factor in time and lost productivity. Cabletron's switched network with SecureFast and its support of the 802.1q VLAN industry standard has all but eliminated this issue [Ref. 34].

With SecureFast in place, most adds, moves and changes can take place automatically, with others being performed quickly through a simple point-and-click interface, even across different protocols. The Network Management Group can also group users based on application, MAC address or geography using the same easy-to-use application.

Support for DHCP is also provided through SecureFast, allowing for the automatic assignment of user IP addresses. This is a benefit-rich feature for many organizations, not just Cabletron, with more and more workforces requiring

Web access to perform their jobs. (In the past, duplicate addresses have often resulted in lengthy delays and lost production not only for those users affected by the error but the MIS staff that has to fix the problem and re-assign new addresses.)

"We used to perform an average of four to five moves a day before the switch over," notes Smart.

And today, Cabletron is still growing at the same pace so we have the same number of moves. But what used to take at least an hour or two before--including tireless visits to the wiring closet, and re-configuring router codes-- now takes less than five minutes each [Ref. 34].

If you put that in man hours and what it could cost a company to pay someone to do all this re-configuration, it's clear that Cabletron is getting much more done in less time. Needless to say, the cost savings are dramatic.

As part of Cabletron's migration to a switch-based network, the company also continued to rely on its SPECTRUM enterprise management platform to oversee the network. This platform is critical in a switched environment because it pinpoints faults to ensure a better running network top to bottom. In addition, SPECTRUM also has new features like Web-based management applications and systems management applications to efficiently streamline the time it takes to oversee an enterprise network. Even if a Cabletron network manager is on the road, he can resolve a problem within a few minutes. Considering that every minute of downtime can cost a company upwards of a \$1,000 or more, this also goes to great lengths to reduce costs [Ref. 34].

Because network security is such an important issue these days, SecureFast provides MIS and business managers with policy-based management tools. "With SecureFast, you can secure certain segments on the network from other users," Dunigan explained.

SecureFast will allow us to determine if a user has the rights to traverse a particular segment of the network. If the user does not have privileges, he will automatically be denied access to, for example, certain file servers on a particular network segment. You can deploy as much security in a switch-based network as you want, and still not impact the overall performance [Ref. 34].

For business managers at Cabletron who have been concerned with how networked resources are being used, they now have a simple way to maintain security, not to mention their own peace of mind.

What about those hidden costs associated with networked users who may be abusing the network, even within their assigned privileges? SecureFast helps with a Call Accounting application which stores information for each connection on the network, including the duration of a call, the type of service (e-mail, video conferencing, etc.), amount of bandwidth it consumed, when it took place and the identity of the connected users. You can liken it to the same service provided by the telephone company [Ref. 34].

With Call Accounting, this detailed information can be made available to whomever needs it within Cabletron: accounting, personnel, department heads, just about anyone. The point-and-click SecureFast application even allows users to format and export the data to a spreadsheet or billing application. The costs associated with the day-to-day operation of the enterprise are provided in plain black-and-white. Now that's cost justification.

As mentioned earlier, the way many enterprise networks operate there is too much time spent with daily operational tasks. When the network grows, this often requires adding more and more support staff just to keep up with this constant workload. The ongoing hiring of highly paid staff does nothing but significantly drive up the cost of ownership.

Cabletron, on the other hand, has grown exponentially as a company, going from 3,000 to more than 7,000 employees worldwide, but the network management staff has only increased to 13 people in all. Just 13 people to manage a worldwide enterprise network with over 20,000 nodes [Ref. 34].

2. The Multiple-Site Solution

The work completed at Cabletron's campus serves as an ideal model for single-site corporate networks. However, like most organizations, Cabletron is not merely a local entity with a need for a single pipe to the outside world. It is a global corporation with more than 200 offices all over the world. And with an enterprise that large, the costs associated with running it can be staggering [Ref. 34]. Especially when you factor in the use of voice and video services being piped into the enterprise.

Since this extends outside the LAN environment and requires the expertise of voice and data communication leaders, Cabletron has looked to team up with other vendors to bring a best-of-breed approach to the cost-effective consolidation of voice, video and data resources.

By pooling talent and technology with the likes of Nortel and MCI, Cabletron can expand its corporate enterprise to build a robust, all-

encompassing network that works to reduce costs as it improves performance and productivity company-wide, see Figure 28.

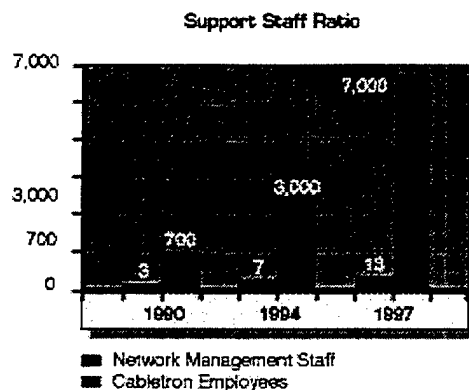


Figure 28. Support Staff Ratio. [Ref. 34]

Nortel and Cabletron share very complementary visions for building next-generation switched networks. Through Nortel's leadership in voice technology and Cabletron's role in data networking, the two companies are able to deliver an integrated solution, capable of supporting critical business applications such as voice, multicast, video on demand, and multimedia. Nortel refer to this consolidated enterprise as a "Power Network"-an architecture created by Nortel to deliver strategic business solutions that satisfy customer requirements for implementing new data and voice services [Ref. 34]. Cabletron provides a similar foundation for enabling enhanced data services across an ever-growing switch-based networking infrastructure. By partnering in technology and business focus, Cabletron and Nortel are uniquely positioned to deliver real business solutions to customers and provide a seamless enterprise that is scalable, simple to implement and maintain, highly distributed and flexible, and able to accommodate growing application demands [Ref. 34].

Cabletron put this partnership right to work in early 1997 by implementing a nationwide LAN/WAN environment, see Figure 29. Leveraging each company's strength in ATM technology, Cabletron has combined its SmartSwitch 9000 (formerly the MMAC-Plus) and SmartSwitch 6000 for LAN connectivity, with Nortel's Meridian 1 Communications System switch for voice connectivity and Magellan Passport ATM enterprise network switch for multimedia transport through the WAN. All of the devices are managed by Cabletron's SPECTRUM Enterprise Manager [Ref. 34].

This "Power Network" has unified over 200 offices around the world comprised of approximately 20,000 nodes. Now the existing voice network and LAN services are located on the same ATM backbone. By using Cabletron's SecureFast architecture, Cabletron employees can receive the same Quality of Service (QoS) in the LAN and WAN that customers have come to expect from their phone service.

According to Smart, this new solution has already resulted in significant cost-savings. "Not only have we reduced complexity by consolidating our voice, video and data on one system, but we've been able to lower WAN costs by about 40 percent [Ref. 34]." The future is even more promising with this new solution.

Cabletron plans to merge the LAN and WAN concepts into one paradigm that will stretch across the world see Figure 29 and Figure 30. This new system will simultaneously support all of Cabletron's necessary business applications including multimedia, desktop computer telephony integration (CTI), videoconferencing, video on demand, and more. Key goals include [Ref. 34]:

- Savings of approximately \$27,000 per month over former divided network operations
- Cost per user actually decreases as new users are added
- On-net video links eliminate bandwidth costs
- Dynamic bandwidth allocation for more efficient voice, video and data transmission

While reducing costs, this consolidated ATM network will provide the following benefits [Ref. 34]:

- Provide more accurate information to customers and employees
- Provide faster service to customers

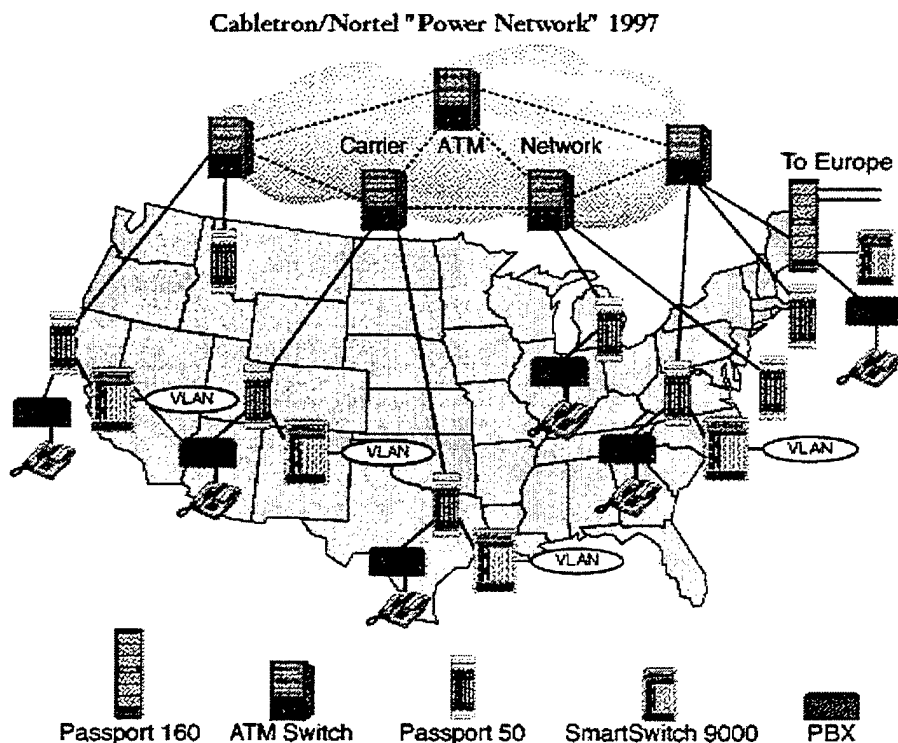


Figure 29. "Power Network 1997". [Ref. 34]

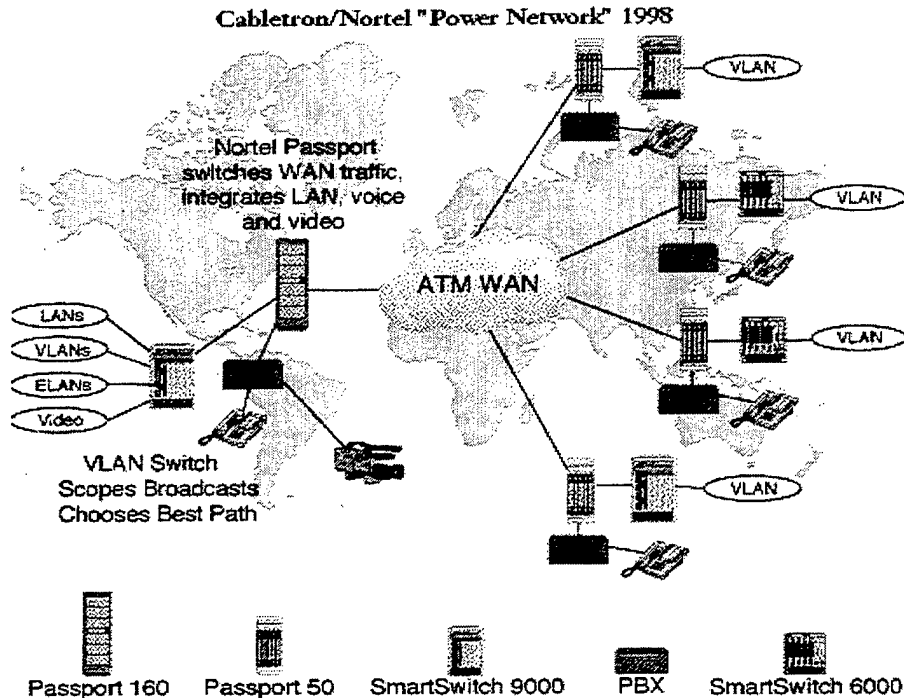


Figure 30. "Power Network 1998". [Ref. 34]

- Savings on telecommunications usage can be passed on to customer or other areas to improve the business

H. A PLAN TO FOLLOW

Now that Cabletron's new enterprise is in place, the Network Management group believes customers will want to emulate this switched-based solution.

Customers who currently have routers at the center of their networks are prime for a network migration similar to Cabletron's. Their networks, like the earlier referenced router-based network, are probably pushing the routers' CPU capacity. Their routers will eventually choke on all the traffic like ours did, Dickson explained. They don't have to get rid of the router; they simply need to re-deploy it and put more of the traffic burden on the switches which can handle the load much more efficiently [Ref. 34].

Smart added,

Software development houses for example, with their massive data transfer needs, demand the bandwidth that our network prototype can provide. Now they can look at our new infrastructure and see

firsthand that by flattening out the network with switch-based products, you remove the administrative nightmare of figuring out subnets [Ref. 34].

Additionally, it removes the tediousness of constant adds, moves and changes. Within a switch-based network, MIS departments no longer have to decide what subnet to put new employees on, or what servers specific groups need to get to, or even what hub to cross-patch employees to when there are ten hubs in one wiring closet.

Companies that can benefit from a switched-based network are not limited to traditional networking industry businesses either. Investment houses, for example, need high-speed backbone and redundancy features in their networks to be able to trade on the fast-paced, bandwidth-hungry stock market. An important point to remember is that regardless of where a customer is with their network migration--whether they're starting from the ground up or simply looking to increase performance within a few departments--they have the flexibility to build the enterprise according to their specific needs and at their own pace.

If the company is only concerned about improving the data throughput of the network, a Fast Ethernet board in the SmartSwitch 9000 may fit the bill. If the company wants to consolidate voice and data in a regional office for a lower bottom line, the SmartSwitch 6000/Passport solution will suffice. And if a router is required to handle a significant amount of Layer 3 and 4 processing, a new hardware-based switch/router device could be the answer. (Unlike traditional software-driven routers, these innovative multi-layer, switch/router devices do all

processing through dedicated ASIC engines, resulting in Gigabit throughput that's nearly 15 times faster at just a fraction of the cost.) [Ref. 33]

Either way, customers who place a substantial investment to follow Cabletron (and its strategic partners) in this network direction will have confidence knowing that it was done for the same reasons: to improve overall network performance and to lower ownership costs.

Dunigan best summarized Cabletron's migration path as putting network control back into the hands of business managers.

Because it is the business managers who are essentially responsible for operating the company, they should have the final say on how the network is utilized. As for myself and other network management administrators, because of our new network infrastructure we have moved from the role of constantly fighting fires to truly managing and enhancing overall network performance [Ref. 34].

VI. CONCLUSION: HIGH-PERFORMANCE NETWORKING INITIATIVES

A. NEXT GENERATION INTERNET NETWORK POLICY APPLICATIONS AT NASA

We must build the second generation of the Internet so that our leading universities and national laboratories can communicate in speeds 1,000 times faster than today, to develop new medical treatments, new sources of energy, new ways of working together. But we cannot stop there. As the Internet becomes our new town square, a computer in every home -- a teacher of all subjects, a connection to all cultures -- this will no longer be a dream, but a necessity. And over the next decade, that must be our goal.

*President Clinton
1997 State of the Union Address*

On September 10, 1997 Rep. F. James Sensenbrenner, Jr. (R-WI), chairman of the House Science Committee, upbraided representatives of the Clinton Administration for their slowness in putting together a detailed plan for implementing the proposed Next Generation Internet (NGI).

In his fiscal year 1998 budget request, the President asked for \$100 million dollars to fund the NGI initiative. The funding was to come out of the budgets of five different federal agencies, primarily the Defense Advanced Research Projects Agency (DARPA, initiator of the original Internet), the Department of Energy (DOE), the National Science Foundation, and NASA. The other contributors are the National Institute of Standards and Technology, and the National Library of Medicine/National Institutes of Health.

The President's original proposal first ran into trouble with Congress in Spring 1998. Despite congressional support for the principles underlying the NGI

proposal, many committees withheld funding for the initiative at that time in order to secure a more detailed implementation plan from the Administration. Many members also expressed concern that the original plan provided only a few select locations with enhanced connectivity. This would have meant that some areas of the country, like Alaska and parts of the Midwest, would be left behind as infrastructural improvements were made. Some also questioned the role of the federal government in improving the Internet, asserting that taxpayer dollars should not be spent on something that industry would do anyway.

In order to more thoroughly address congressional concerns, the Administration revised the NGI proposal, releasing a new draft implementation plan in July 1998. The new plan identifies NGI as a research initiative, rather than a deployment initiative, more clearly than the original proposal. Its first stated goal is to conduct experimental research to develop advanced network technologies, with DARPA as the leading federal agency. Goal number two is to implement a high-speed "network fabric" which will provide the means to test new technologies. This fabric will provide connectivity about 100 times faster than the current Internet to at least 100 universities and federal research sites. Goal three is to come up with "revolutionary applications," new ways of using the technologies that will emerge from the NGI research. Hopes are high for this aspect of the initiative which has the potential to allow unimaginable means of communicating, collaborating, and creating. Based on these three goals. The Next Generation Internet will attempt to [Ref. 38]:

- Promote experimentation with the next generation of networking technologies.

- Develop a next generation network testbed to connect universities and Federal research institutions at rates that demonstrate new networking technologies and support future research.
- Demonstrate new applications that support important national goals and missions such as scientific research, national security, distance education, environmental monitoring and health care.

To achieve these goals, NGI will be built on the base of current R&D activities and programs in the participating Federal agencies. Furthermore, it will call on substantial matching funds from its private sector partners and collaborate with academia. These three goals are explained further in section 3.

Since the language of the new NGI Draft Implementation Plan is focused on research and development, the door is left wide open for private interests to actually implement the technologies that emerge. This may help assuage the concerns of members of Congress who detected possible "corporate welfare" in the initial version of the plan. In addition, the revised plan's Administration proponents are specifically addressing the concerns of members of Congress from areas that would have a tough time participating in NGI.

The Presidential Advisory Committee [on NGI] is very concerned about this problem because we believe that networking should bring all parts of the Nation together rather than amplify any geographical disadvantages,

stated Dr. John Gibbons, director of the Office of Science and Technology Policy on September 10, 1998. Dr. Gibbons went on to explain that some areas, like Alaska, would require far more money than NGI had to offer to bring them up to full speed for the test network. However, the federal agencies involved in NGI are being encouraged to provide increased funding to disadvantaged locations.

However, despite these efforts to assuage congressional concerns, it may be too late to save funding for NGI in 1998 according to Rep. Sensenbrenner. "The trains going to leave the station," he said at a hearing on September 10, 1998.

1. The Results of the Game Plan: NASA Participation in NGI

The National Aeronautics and Space Administration, NASA, is participating in the Next Generation Internet (NGI) initiative, a multiagency effort that also includes the Departments of Defense, Energy, and Commerce, and the National Science Foundation. NASA's NGI focus is on prototyping NGI applications [Ref. 47].

a. Technologies

NASA will deploy an appropriate suite of advanced networking services to enable high performance applications. NASA-sponsored research will focus on important issues such as network performance measurement, network interoperability, quality of service and network security. NASA will fund and manage research in advanced network technologies that are richer in features, higher in performance, and deliverable at a reasonable cost [Ref. 38]. For example, they will enable real-time networking, group collaborations, remote access to the network, and a seamless interface for space-to-ground communications.

NASA will continue to be an early adopter of emerging networking technologies that chart a course for a robust, scaleable, shared infrastructure supporting lead users from NASA, other government agencies, and the research

community, as well as large numbers of ordinary commercial users [Ref. 47].

NASA's program goal relevant to NGI's goal 1 is to sponsor R&D in new networking technologies and services in support of the high performance applications requirements. NASA will partner with industry and academia on R&D in internetworking technologies to achieve an interoperable high performance network testbed. By doing so, NASA will deliver advanced networking technologies to the aerospace community and ultimately to the public.

b. Testbeds

NASA will provide both a high performance network application testbed and a network research testbed for the NASA community and its partners. NREN is the NASA NGI testbed and will enable NASA to focus on delivering a leading-edge application environment to its community. Therefore, NASA will [Ref. 38]:

- Enable next-generation application demonstrations across the network. Internetwork with other Federal agencies and academic and industry partners at both the IP and ATM service level.
- Deploy advanced networking services such as IPv6, multicast, QoS, security and network management tools.
- The research effort supports such advanced capabilities as remote interactive graphics, international digital libraries, and network-based high-definition displays for science, manufacturing and education.

Among the features to be packed into the NGI program will be new switching systems, network protocols, high-speed interconnections to workstations and supercomputers, as well as new forms of interconnection and hybrid networking to reach remote and mobile users. NGI will highly leverage

industry developments and any ongoing Federal research to provide a hybrid networking demonstration platform. NASA will leverage its experience in high-speed satellite data communications from the Advanced Communications Technology Satellite program and attempt to make use of existing NASA satellite resources as well as seeking out satellite services from commercial sources. The high-speed satellite links mentioned above could provide a means of connecting international testbeds to the NGI (e.g., GIBN, the Global Interoperability Broadband Network) [Ref. 39].

Managing the dynamics of these activities will be a major challenge, but the payoff for success will be enormous in terms of national capabilities, research productivity, and new commercial products and services [Ref. 39]. This initiative is important to NASA because NASA missions require the interconnection and integration of its unique resources that include user facilities, databases, and supercomputers, as well as geographically distributed researchers and scientists located at universities, federal research institutions and in industry [Ref. 38].

The NASA-funded research and engineering community consists of over 180 universities, 30 industry partners, and 5 science and research centers where NASA has significant research programs in place [Ref. 47]. The importance to NASA of coupling resources at different sites to solve critical problems is underlined by the research challenges in the fields of advanced aerospace design, Earth sciences, astrobiology, astrophysics, telemedicine, multicast network technology and space exploration [Ref. 47].

NASA has already embarked on a number of applications which will require the network technology acceleration of the Next Generation Internet to be successful. Sample revolutionary applications include [Ref. 38]:

- Advanced Aerospace Design and Test Tools - NASA wind tunnels on-line, virtual flight simulation laboratories on-line.
- Telemedicine - Interactive consultations, remote protocols and procedures modeling distant health care delivery in space.
- Earth Sciences - Advanced science investigations for Mission to Planet Earth.
- Astrobiology - Remote scientific analysis of Martian rock, virtual aerospace environments for distributed collaborations.
- Astrophysics - Remote operations of space telescopes located in isolated areas such as the Keck Observatory in Hawaii.
- Space Exploration - Remote interactive visualizations for command and control of robotic explorers such as the Mars Pathfinder.

Unfortunately, the current Internet is too unreliable, too primitive, too geographically limited, and has too low a capacity to provide strong support for these requirements. NASA and other agencies have been working on technologies in concert with private industry to explore new networking technologies. The NGI will address not only accessible but also remote sites and rural states. NASA experiments are anticipated to assist research in reaching beyond the current Internet infrastructure to accelerate technology development and deployment to remote locations. This initiative provides the critical mass and leverage to unite this work and bring it to rapid fruition [Ref. 39].

2. Naval Postgraduate School Contribution: Server and Agent Based Active Management (SAAM) Architecture

Researchers at the Naval Postgraduate School are developing a SAAM system for the NGI. The SAAM project is currently sponsored by DARPA and NASA Research and Education Network (NREN) under the NGI initiative. One major objective of the NGI is to support all types of data using a single network. This integrated services requirement poses a significant new challenge to network management: namely, Quality of Service (QoS) path management. Specifically in the NGI, the capacity of each link will be shared by a set of logical service pipes, each of which provides a particular level of packet performance measured by a set of QoS parameters [Ref. 40]. Typical QoS parameters include the bound on packet delay and the rate of packet loss. For a data flow (e.g., a video flow) that requires end-to-end QoS guarantees to its packets, the source will invoke a resource reservation protocol such as RSVP [Ref. 40] to establish a QoS path to the destination. The path is composed of a sequence of service pipes whose composite QoS meets the flow's requirement. In summary, in addition to maintaining connectivity, the NGI must dynamically allocate and maintain QoS paths. SAAM servers and the routers will use a real-time transport protocol for data and agent communications. Therefore in SAAM, most management and control tasks will be performed by the servers in an automated and timely fashion. Only a small number of planning tasks will require human interaction, and in such cases, the management station will need to communicate with a high level SAAM server most of the time. SAAM will also have built-in mechanisms to interact with the reservation protocol and provide it useful path

information when requested. SAAM will also support Diff-serv and MPLS in addition to per-flow reservation for real-time data [Ref. 40].

3. Other High-Performance Networking Initiatives

The two major U.S. initiatives in high-performance networking are the Next Generation Internet (NGI), and Internet2 (I2), [Ref. 46]. Internet2 is a collection of 135 universities and businesses, including Stanford, the University of California at Berkeley, Harvard, Cornell, Yale and the University of Virginia. They are trying to create a sort of virtual university for students and professors to access books from libraries thousands of miles apart, to take classes at other campuses and to collaborate on research projects.

Formed last October, its ultimate goal is to connect campuses at speeds so fast that a 30-volume encyclopedia could be transmitted in less than a second [Ref. 39].

Although separate projects, NGI and Internet2 share many things in common. Both have limited links to the commercial Web, and both will depend on the Internet for e-mail, low-level research and other day-to-day uses; both are being built simultaneously; and some elements of Internet2 are being integrated into NGI [Ref. 47].

Although the new Nets at first will be devoted to research, both may become available for electronic commerce -- which quickly is becoming a major focus of the Internet.

The National Science Foundation (NSF) Very High-Speed Backbone Network Service (vBNS) [Ref. 45] plays a key role in both the NGI and I2. The

goals and strategies of NGI and I2 have striking similarities:

- Both seek to develop mechanisms for differentiating among the services that the Internet provides to a variety of application types; this ability to differentiate is often called Quality of Service or QoS
- Both have developed private industry-government-academia partnerships
- Both seek to transfer their results to widespread practice in the commercial Internet as soon as possible (as neither initiative seeks to develop and maintain its own networks for long periods of time)

Several other countries have commissioned similar initiatives with similar objectives. Canada has an initiative called the Canadian Network for the Advancement of Research, Industry, and Education (CANARIE) [Ref. 46]. CANARIE is also an industry-government-academia partnership. CANARIE commissioned CA*net II to develop advanced Internet capabilities and to transfer them into the private sector. Singapore carved out of its commercial backbone the Singapore Internet Next Generation Advanced Research and Education Network (SINGAREN). Taiwan has TANet. [Ref. 46] France is readying its Renater-2. NORDUnet, the academic backbone network for the Nordic countries, is approaching a NORDUnet-2 initiative. Germany is commissioning an advanced high-performance network. The Asia-Pacific Advanced Network forum (APAN) is a cooperation of several countries, including Japan, Korea, Singapore, and Australia (founding members), joined by Hong Kong, Indonesia, and Thailand.

The vBNS, CA*net II, and SINGAREN are already connected to STAR TAP. The US National Science Foundation (NSF) sponsored Science,

Technology, and Research Transit Access Point (STAR TAP) initiative provides an international exchange point for high-performance networked applications. STAR TAP will be one of the exchange points for the Next-Generation Internet (NGI). Also, STAR TAP is housed in the same facility as one of the emerging Internet 2 gigapops.

TANet will soon connect, and APAN is seeking a way to connect its members' advanced networks. Several European advanced networks, mentioned above, are pursuing connections. In addition to NSF, other U.S. government agencies that participate in NGI are connected at the STAR TAP; these include the Department of Energy's Energy Sciences Network (ESnet) [Ref. 46] and the National Aeronautics and Space Administration's NASA Research and Education Network (NREN) [Ref. 44].

4. Networking Policy Considerations For Next-Generation Applications

Ideally, application requirements drive advanced network design. A number of advanced research networks are being established around the world to help develop the NGI network testbeds and applications [Ref. 45]. These networks support trial activities that are not practical using the regular Internet. These activities include QoS, IPv6, some forms of multicasting, etc. More importantly, these advanced networks support high-performance applications that will not work properly, if at all, on the big "I" Internet--applications such as medical imaging, tele-immersion, collaboratoriums, distributed genome sequencing, etc.

It has been long recognized that most high-performance applications development occurs at leading university and research institutes. As in the early days of the Internet, the research community pushes the envelope in terms of high-performance applications. Funding and infrastructure support enables not only these applications, but basic research as well. A "tertiary" benefit of applications development is that network requirements and design are driven by user "needs" rather than network designers' perceptions of what is required. Satisfying the real needs of researchers will allow the development of applications and services that can easily migrate into the commercial sector [Ref. 45].

A classic example of this "tertiary" benefit is the development of the World Wide Web. It was originally developed to serve the needs of high energy researchers. It is expected that today's next-generation networks will lead to similar application developments and spin-offs [Ref. 45].

According to NREN project researchers at the NASA AMES Research Center, the next-generation of Internet networks being deployed worldwide generally have restrictive Authorized User Policies (AUP) that limit connectivity to a subset of institutions carrying out high-performance meritorious applications [Ref. 47]. Currently, in order to interconnect "cooperating" NGI network testbeds, network operators enter into peering agreements where they identify specific institutions that are permitted to communicate with each other across a common interconnect point, like STAR TAP. The routes for these institutions are then

advertised by each respective network to their member institutions.

The use of community attributes is, however, only a partial solution and only works properly if the foreign network has implemented an "explicit" routing architecture.

a. *Specific Policies*

- Quality of Service
- Security
- Policy Architecture: Large Network Requirements

(1) QoS. There is wide-spread agreement that the Internet needs some sort of "Quality of Service" (QOS) capability [Ref. 41]. Congestion associated with the phenomenal growth of the Internet and World Wide Web has made the Internet essentially useless for "production" or mission-critical work. QOS offers the potential of ensuring bandwidth availability for critical needs while still offering the current, best-effort service for non-critical traffic. In the longer term, economic theory suggests that some sort of service differentiation (and associated price differences) is necessary to make the Internet a viable, self-sustaining, commercial enterprise. For long term sustainability, a competitive entity needs money both to support its operating costs and to fund its expansion [Ref. 41]. A collection of such entities need the spectrum of services to create "market niches" that help prevent the largest entity from using its size to create a monopoly. QOS could provide the different services and prices needed to meet these needs. In this research the author found that providing QOS, a QOS that "fits" the Internet, seems to be the single most important, and most difficult,

challenge to its long-term survival. However, while there is agreement on the need for QOS, there is none on how QOS should be added to the Internet/IP communication model. Currently, NASA AMES NCI network researchers are experimenting with a Differential Services QoS model emphasizing different queueing mechanisms [Ref. 47].

(2) Security. There is no doubt that a next generation Internet will need to have security designed in from the beginning. It should be noted, though, that for the most part, the problem today is not with the design of the protocols or the net. The Next Generation Internet should be designed so that cryptography is ubiquitous, largely transparent, and used universally where appropriate for authentication and/or privacy [Ref. 42]. Even at a more abstract level, it is not clear that today's cryptographic devices are adequate. If nothing else, encryption, hashing, and authentication algorithms are quite expensive. Given the limited line speeds in today's long-haul nets, client hosts can easily keep up; as line speeds increase, the load will become more of a burden, especially for servers. But cryptography cannot solve "the" Internet security problem. An analysis of the CERT advisories for the last 18 months shows that two or three, out of about 30, would be moot if cryptography were universally deployed. Most of the problems are due to buggy software, bad administration, or both. There are several possible approaches. One approach is a Next Generation Orange Book. While not a bad idea, the philosophical approach should be re-visited [Ref. 42]. The Orange Book is predicated on the idea that there is one central security model and one security kernel whose job is to

enforce the model. That approach doesn't scale well to today's Internet. Consider the case of a service bureau machine running Web servers on behalf of several mutually suspicious companies. The isolation between companies -- an administratively-decreed structure -- can probably be accommodated within the constraints of the traditional model. The network interface, however, is open to all; there is no barrier to sending anything out, or receiving anything. The individual server programs have their own trust structure. Credit cards numbers - - data that is only a few bytes long -- must be rigorously protected [Ref. 42]. But the card number verification process may require that the numbers be sent to a bank. Some merchants even store credit card numbers in a long-term database, so that they need not be entered each time. How can this fine-grained protection be implemented? Another consequence of buggy code is that firewalls will continue to be needed. However, in an interview NASA AMES NREN network engineers stressed that their nature will change [Ref. 47]. First, the big corporate design won't work. Better policy controls are needed on what can pass through a firewall. This may imply more application proxies, but they have to be much faster and much more flexible. Second, there are some firewall-hostile protocols out there. To give just one example, TCP is much more easily passed by a firewall, because header bits can differentiate, in a context-independent fashion, between incoming and outgoing calls. This doesn't work with UDP. But in reality, despite its datagram nature, most of the useful applications built on top of UDP use a query-response protocol. A different design would have permitted easy firewall filtering of such messages [Ref. 42].

(3) Policy Architecture. Scaling a large network is a principal concern, especially at high speeds, and an issue which needs to be carefully examined in the conceptual design and planning stages [Ref. 43]. Striking the right hierarchical balance can make the difference between a high-performance network and one which is on the verge of congestion collapse. This section discusses the architectural concepts of applying policy control in a large network to minimize the performance overhead associated with different mechanisms that provide traffic and route filtering, bandwidth control (rate-limiting), routing policy, congestion management, and differentiated classes of service (CoS).

- **Architectural Concepts** - The key to scaling very large networks is to maintain strict levels of hierarchy, commonly referred to as "core, distribution, and access" levels, which limits the degree of meshing between nodes. The "core" portion of the hierarchy is generally considered the central portion, or backbone, of the network; the "access" level represents the outer-most portion of the hierarchy, and the "distribution" level represents the aggregation points and transit portions which lie between the core and access levels of the hierarchy [Figure 31].

A high degree of meshing affects the ability to maintain stability in the routing system and degrades overall performance as the network grows larger [Ref. 43]. A distinction should be made between full-meshing and partial meshing; it stands to reason that partially-meshed networks scale better than fully-meshed networks, and for redundancy considerations, a partially-meshed network has acceptable scaling properties (in most instances).

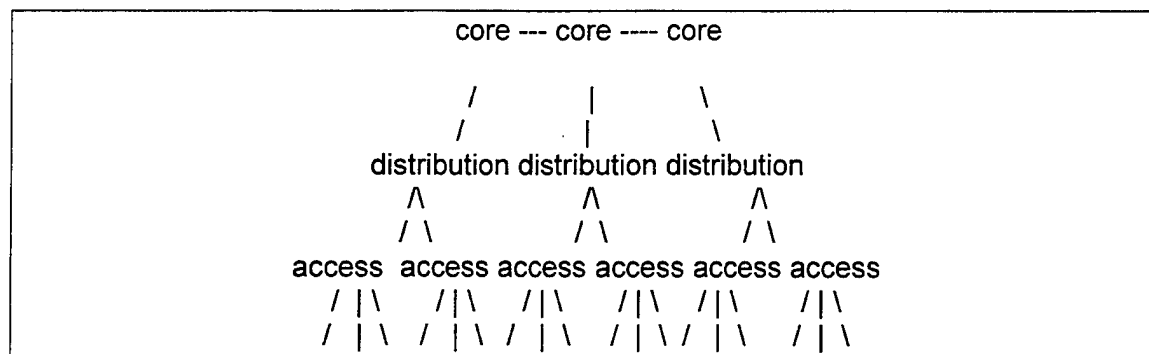


Figure 31. Access Levels of Hierarchy. [Ref. 43]

It is important to understand, however, the impact of applying policy at various locations with the network hierarchy. Depending upon where these policies are applied, it can have varying effects on the performance of the network. It is commonly accepted that the network core should never become the point of network congestion, but in reality, this is not always the case [Ref. 43]. The purpose of the network core is to forward traffic as quickly as possible to its destination, therefore any type of policy which may affect forwarding performance should not be implemented in the network core. Types of policies which may fall into this category are traffic and route filtering. It is worthwhile to consider "pushing" these types of policy implementations to the lower portions of the hierarchy to avoid performance degradations which affect the entire network, instead of a smaller subset of the overall topology. If traffic congestion is a concern in the distribution and core levels of the network hierarchy, then one should consider the implementation of Random Early Detection (RED) as a congestion management mechanism in these portions of the network topology [Ref. 43]. Implementation of policy lower in the hierarchy has a nominal impact

on the overall performance of the network for several reasons. Some mechanisms, such as traffic filtering, have less of an impact on forwarding performance on lower speed lines. Since the speed of the inter-node connectivity generally gets faster as one goes higher in the network hierarchy, the impact of implementing policy in the higher levels of the network hierarchy increases. The same principal holds true for traffic accounting, access control, bandwidth management, preference routing, and other types of policy implementations. These mechanisms are more appropriately applied to nodes which reside in the lower portions of the network hierarchy, where processor and memory budgets are less critical. Another compelling reason to push policy out towards the network periphery is to maintain stability in the core network. Since high speed traffic forwarding is usually found in the network core, and conversely, lower speed forwarding is generally found lower in the network hierarchy, there can be varying degrees of performance impact depending on where these types of policies are implemented. In most cases, there is a much larger percentage of traffic which transits the network core than transits any one particular access point, so implementing policy in the network core has a higher degree of impact on a larger percentage of the traffic [Ref. 43]. By the same token, any adverse performance impact due to such policy implementations in the higher levels of the network hierarchy has a broader impact on a larger percentage of the overall network. Policy implementation in large networks should be done at the lowest portions of the hierarchy as possible, in order to avoid performance degradations which impact the entire network [Ref. 43].

B. RECOMMENDATIONS/SUGGESTIONS FOR FURTHER STUDIES AND RESEARCH TOPICS

1. Gemini, an Example of the Need for Research Collaboration Linkages

Today's high-performance applications have increasing requirements to establish network topologies that follow research collaboration linkages rather than network topologies [Ref. 44]. For example, there is a joint project among Canada, the NASA Goddard Space Flight Center (GSFC), United Kingdom, Chile, Argentina, and Brazil called Gemini, for two very powerful next-generation "twin" telescopes to be built in Chile and Hawaii. The telescopes will use atmosphere-correcting lenses, which should result in images almost as good as those produced from the Hubble space telescope. Also, these telescopes will use electronic imaging systems rather than film. Images could be transmitted over high-performance networks to a researcher's desk. In time, it should be possible for researchers to operate the telescope remotely in real time [Ref. 44].

These images cannot use any lossy compression technique, as very faint star images may be lost in the process; hence, image files will be several megabytes, if not gigabytes, in size. These images will be exchanged back and forth among a small set of researchers located primarily in the participating countries. This application requires high-performance connectivity among sites, yet currently each site is connected to separate research networks [Ref. 44].

With the existing interconnection agreement, collaboration linkages among researchers are driven by the somewhat arbitrary designation of approved

institutions in respective networks [Ref. 45]. In an ideal world, it would be advantageous to set up "wide-area virtual high-bandwidth networks" among collaborating institutions. These virtual high-performance networks would span a number of underlying research networks. The analogy that is commonly used is a set of overlapping Venn diagrams of virtual private networks (VPNs). These are somewhat different in concept from today's commercial I-VPNs in that an institution can be a member of a number of different communities of interest. More important, users are not isolated from the Internet as they would be with I-VPNs. In fact, the Internet could be considered the global "community of interest" of which every Internet user is a member [Ref. 45].

Using the Gemini project as an example, Canada's University of Toronto, the U.S.'s NASA GSFC, Chile's National University of Chile, and select institutions in the United Kingdom, Argentina, and Brazil could be part of one virtual private network that is dedicated to one specific collaborative application [Ref. 46]. Using what are commonly referred to as "explicit routing" technologies, such as MPLS, these communities of interest could be assigned their own QoS mechanisms where not only the advertised routes but other parameters, such as bandwidth and delay, could also be configured dynamically.

Another important feature of the "communities of interest" concept is that they are layer 2 independent. There is strong likelihood that next-generation networks, including SONET, Giga-ethernet, Optical IP, and ATM, will use a number of layer 2 transport services. A community-of-interest or VPN strategy that is layer 2 independent is very attractive to many network operators [Ref. 47].

The need for virtual high-performance networks will be increasingly important for funding agencies as well. In the next few years, it will be increasingly difficult for funding agencies to continue to fund "generic" backbone networks for the research community, even if those networks have restrictive AUPs [Ref. 46]. On the other hand, it is recognized that many advanced institutions and applications need high-performance networks to carry out their fundamental missions. Virtual networks that are clearly identified with a specific research objective or even a specific application will be much easier to support from public funds [Ref. 45].

2. IPV6

The Internet Engineering Task Force (IETF) adopted Internet Protocol, Version 6 (IPv6) [RFC 1883] as a replacement for Internet Protocol, Version 4 (IPv4). Research could be conducted to provide a brief technical assessment of IPv6 and specifically cover such major aspects of IPv6 as routing and addressing, address autoconfiguration, neighbor discovery, and transitioning networks from IPv4 to IPv6. Other aspects of IPv6, such as support for mobility, support for resource reservations, and security should also be addressed in the research [Ref. 46].

The ability to sustain continuous and uninterrupted growth of the Internet could be viewed as the major driving factor behind IPv6. IP address space depletion and the Internet routing system overload are some of the major obstacles that could preclude the growth of the Internet.

One of the fundamental requirements of the current IP routing and addressing architecture is the requirement that within a single internet, such as the Internet, IP unicast addresses must be unambiguous (unique). In other words, within an internet, only one node could have a given IP unicast address. Combined with the fact that an IPv4 address is 32 bits wide, this means that the theoretical upper bound on the size of the Internet is 2^{32} nodes, consisting of hosts and routers. By extending the size of the address field in the network layer header from 32 to 128 bits, IPv6 raised this theoretical limit to 2^{128} nodes. Therefore, IPv6 could solve the IP address space depletion problem for the foreseeable future [Ref. 46].

Although IPv6 significantly increases the total size of the available IP address space, the question remains: Is this sufficient to enable the continuous growth of the Internet? Within the current IP routing and addressing architecture, IP addresses must be unambiguous, but this is not sufficient to guarantee IP-level reachability within an internet [Ref. 46]. Within the current architecture, the primary role of IP addresses is not just to enumerate all the nodes (hosts and routers) within an internet, but to provide routing (IP-level reachability) to all the nodes within the internet. Therefore, in order to support an uninterrupted growth of the Internet while maintaining the current IP routing and addressing architecture, not only is a larger IP address space needed, but the assignment of addresses must also enable scalable routing [Ref. 46].

LIST OF REFERENCES

1. Charting the Seas of Information Technology: Chaos, The Standish Group International Inc., 1994.
2. See Information Technology Investment: A Government wide Overview (GAO/AIMD-95-208, July 31, 1995).
3. Executive Guide: Improving Mission Performance Through Strategic Information Management and Technology (GAO/AIMD-94-115, May 1994).
4. IT investment is defined as an expenditure of money and/or resources for IT or IT-related products and services involving managerial, technical, and organizational risk for which there are expected benefits to the organization's performance. These benefits are defined as improvements either in efficiency of operations or effectiveness in services (such as reductions in process cycle time or operational costs, increases in speed or quality of customer service, or improvements in productivity).
5. Henderson, Tom, "The Bandwidth Bottleneck--Plan According...", "Windows Magazine, 1 March 1997.
6. Case, J., Fedor, M., Schffstall, M., Davin, J., "A Simple Network Management Protocol (SNMP)", RFC 1157, May 1990.
7. Meyer, K., Erlinger, M., Betser, J., Sunshine, C., Goldszmidt, G., Yemini, Y., "Decentralizing Control and Intelligence in Network Management", Proceedings of International Symposium on Integrated Network Management, May 1995.
8. Leinwand, A., Fang, K., "Network Management: A Practical Perspective", Addison Wesley, 1993.
9. Applegate, L., McFarlan, F., and McKenney, J., Corporate Information Systems Management Text and Cases, 4th edition, Irwin, 1996.
10. NetManage Policy Management Architecture., <http://www.Network.com>.
11. Stamatelopoulos, F., Chiotis, T., Maglaris, B., A Scalable, Platform-Based Architecture for Multiple Domain Network Management".

12. Douglas W. Stevenson. Network Management: What it is and what it isn't; Predictive Network Management call for White Papers. Douglas W. Stevenson. April 1995.
13. Goldszmidt, German, "On Distributed System Management" In Proceedings of the Third IBM/CAS Conference, Toronto Canada, Oct 1993.
14. Case, J., Fedor, M., Schffstall, M., Davin, J., "A Simple Network Management Protocol (SNMP)", RFC 1157, May 1990.
15. Meyer, K., Erlinger, M., Betser, J., Sunshine, C., Goldszmidt, G., Yemini, Y., "Decentralizing Control and Intelligence in Network Management", Proceedings of International Symposium on Integrated Network Management, May 1995.
16. Wack, John, Network Policy Management., <http://www.Netpolicy.com>.
17. Michael, J.B., In the Doctorial Dissertation: A formal process for testing the consistency of composed security policies. George Mason University, Fairfax, Virginia, Spring 1993.
18. CiscoNetworkManagementTechnicalPapers.<http://cco.cisco.com/warp/public/734/capn/technical.shtml>.
19. Erica Roberts. Policy-Based Networking: The New Class System. CPM Net network management. <http://www.cmpnet.com/>.
20. Hyde, Douglas, Web-Based Management., <http://www.3com.com/>.
21. The Next Generation of Network Management., <http://192.156.136.22/technology/>.
22. Bieber M., Vitali F., "Toward Support for Hypermedia on the World Wide Web", *Computer*, Jan, 1997, pp. 62-70.
23. "Issues and Trends: 1998 CSI/FBI Computer Crime and Security Survey," March 4, 1998.
24. High Risk Series: Information Management and Technology (GAO/HR-97-9, February 1997).
25. Financial Audit: 1997 Consolidated Financial Statements of the United States Government (GAO/AIMD-98-127, March 31, 1998).

26. Actually Useful Internet Security Techniques, Larry J. Hughes, Jr.; New Riders Publishing, 1995.
27. Building Internet Firewalls, D. Brent Chapman and Elizabeth D. Zwicky, O'Reilly & Associates, Inc.; 1995.
28. <http://www.checkpoint.com/products/technology/stateful.html>
29. Getting Started with FireWall-1, User Guide, version 3.0, 1997, Check Point Software Technologies, Inc.
30. <http://www.checkpoint.com/products/firewall-1/descriptions/products.html>
31. <http://www.checkpoint.com/products/firewall-1/descriptions/authentication.html>.
32. <http://www.checkpoint.com/products/firewall-1/descriptions/acontrol.html>.
33. Enterprise Management Professional Services, Inc. 1 March, 1997 WWW.emps.com.
34. <http://www.cabletron.com/white-papers/reduce-cost/>.
35. <http://www.igd.fhg.de/~likavec/netman/papers.htm>.
36. <http://www.igd.fhg.de/~likavec/netman/costred.htm>.
37. <http://www.nren.nasa.gov/aboutngi.htm>.
38. <http://www.riacs.edu/research/project.list.html>.
39. <http://www.nisn.nasa.gov/doc.repos/secplan.html>.
40. Xie, Geoffrey G., "SAAM: A Network Management System for the NGI" Department of Computer Science, Naval Postgraduate School <http://www.cs.nps.navy.mil/people/faculty/xie/SAAM/whitepaper.txt>.
41. Jacobson, Van, "Quality of Service for the Next Generation Internet" Network Research Group, Lawrence Berkeley National Laboratory <http://www.cra.org/policy/ngi/papers>.
42. Bellovin, Steve, "Security for the NGI" AT&T Labs research available at <http://www.cra.org/policy/ngi/papers>.
43. Ferguson, Paul, "Policy Architecture in Large Networks" Cisco Systems, Inc. <http://www.cra.org/policy/ngi/papers>.

44. Canadian Network for the Advancement of Research, Industry, and Education, www.canarie.ca B. Davie, Y. Rekhter, and P. Doolan, Switching in IP Networks. Morgan Kaufmann Publishers, San Francisco, CA, 1998.
45. T.A. DeFanti, M.D. Brown and R. Stevens (Guest Editors), "Virtual Reality Over High-Speed Networks," IEEE Computer Graphics & Applications, July 1996, Vol. 16, No. 4, pp. 14-17, 42-84.
46. T.A. DeFanti, I. Foster, M. Papka, R. Stevens, and T. Kuhfuss, "Overview of the I-WAY: Wide Area Visual Supercomputing," International Journal of Supercomputer Applications and High Performance Computing, Vol. 10, No. 2/3, Summer/Fall 1996, pp. 123-131.
47. Marjory Johnson, NASA AMES Senior Scientist, Advanced Computer Center. Personal Interview. 17 June 1997.

BIBLIOGRAPHY

3Com, "Transcend Networking: A Framework for Pervasive Networking," White Paper, 1998.

3Com, "TranscendWare Software--Delivering Policy-Based Networking," URL: www.3com.com-dnsc/600256.html, 11 March 1997.

Advanced Computer Communications, "Frame Relay and Routers," White Paper, URL: <http://www.acc.com/Internet/technology/whitepapers/framerelay.html>, 1998.

Alles, Anthony, "The Next-Generation ATM Switch: From Testbeds to Production Networks," Cisco Systems White Paper, URL: http://www-europe.cisco.com/warp/public/730/General/_wp.htm, 1997.

Anixter, "Enterprise Networks: LANE Update," White Paper, URL: <http://www.anixter.com/1757751.htm>, 1997.

Anixter, "Enterprise Networks: Network Model," White Paper, URL: <http://www.anixter.com>, 1997.

Asanti Tech Note Virtual LAN Discussion," URL: <http://www.asante.com/FAQ/faq.html>, (15 March 1997).

Asbrand, Deborah, "Users Get a Leg Up," LANTimes Online, URL: <http://www.lantimes.com/97/97jun/> 9 June 1997.

Baum, David, "Kodak Develops Bandwidth, Film Company Makes the Move from Shared to Switched LAN'S," LANTimes Online," <http://lantimes.com/96may/> 23 March 97.

Betser, and Bannister, "Decentralized Network Management," <http://www.ito.darpa.mil/Summaries95/A662--Aerospace.html>, (21 March 1999).

Butler, Janet, "Does Chargeback Show Where the Buck Stops?," Software Magazine, v. 12, n. 5, April 1992.

Cisco Systems, "Benefits of Using Cisco IOS," URL: <http://www.cisco.com/warp/public/732/iosben.html>, 29 July 1998.

Cisco Systems, "Cisco IOS Software Features for Differentiated Class of Service for Internetworks," White Paper, URL: http://www.cisco.com/warp/public/732/General/cos_wp.htm, 3 October 1997.

Cisco Systems, "Cisco Multimedia Blueprint Allows Companies to Deploy Networked Applications Today," URL: <http://www.cisco.com/warp/public/146/183>, 23 January 1997.

Cisco Systems, "Fiber Distributed Data Interface," Internetworking Technology Overview, URL: <http://www.cisco.com/univercd/data/doc/cintmet/ito/55773.html>, 1997.

Cisco Systems, "Networked Multimedia Overview," White Paper, URL: <http://www.cisco.com/warp/public/614/19.html>, 17 April 1997.

Cisco Systems, "Token Ring/IEEE 802.5," URL: <http://www.cisco.com/univercd/data/doc/cintimeVito/55031.htm>, 1996.

Cisco Systems, "VLAN Interoperability: VLAN Standardization Via IEEE 802.1", URL: <http://www.cisco.com/warp/public/537/6.html>, (14 July 1997).

CNet Technology, "100Mbps Networking," White Paper, URL: <http://www.cnet.com.tw/support/paper95.html>, 4 August 1997.

Cohen, Jodi, "The Shrinking World of ATM: Cell Technology Getting Squeezed Out of Campus Backbones," Network WorldFusion, 30 June 1997.

Cohen, Jodi, "3Com Primes PoUcy Management for VLANs," Network World, 5 May 1997.

Communications Week, "What's the Best Way to Switch IP?," URL: <http://techweb.com/cw/Web-Links/swtioch.html>, 2 December 1996.

Conrad, James W., Handbook of Communications Systems Management, Third Edition, Auerbach Publications, 1994.

Congestion Control in ATM, URL: <http://www.cne.grnu.edu/modules/atrn/ATMcon.html>, (15 March 1997).

Connected: An Internet Encyclopedia, "RFC 1812 - 5.3.6 Congestion Control," URL: <http://www.dsi.unive.it/Connected/RFC/1812/122.html>, 15 March 1997.

Cooper, Donald R., and Emory, C. William, Business Research Methods, Fifth Edition, Irwin, 1995.

d-Comm, "ATM on a Roll at Last?," URL: <http://www.d-comm.coni/s-bin/sr-read/781>, March 1997.

d-Comm, "Bridging the Networking Gap," URL: <http://www.d-comm.corn/s-bin/sr-read/243>, March 1997.

d-Comm, "Can Network Outsourcing Solve LAN Management Problems?," URL: <http://www.d-comm.com/s-bin/sr-read/440>, (March 1997).

d-Comm, "Fast But Not Furious," URL: <http://www.d-comm.com/s-bin/sr-read/82>, September 1995.

d-Comm, "LANs Switch Away From Routers," URL: <http://www.d-comm.coni/s-bin/sr-read/76>, March 1997.

d-Comm, "SNW and the Future of Network Management," URL: <http://www.d-comm.com/s-bin/sr-read/437>, (March 1997).

d-Comm, "Switching on Cisco," URL: <http://www.dcomm.com/s-bin/sr-read/730>, (February 1997).

Darling, Charles C., "Ethernet Backbone Switches: Road to ATM," Dalamatation, 1 January 1996.

Darling, Charles C., "Routers Can Save Your WAN Dollars: Sophisticated Tricks Help Keep Your Enterprise Pipes Flowing at Peak Efficiency," Dalamatation, URL: www.datamation.com/PlugIn/issues/July1/html, 1 July 1997.

Darrow, Barbara, "3Coin Touts Advantages of Picking Up Pace," Computer Reseller News, 14 November 1994.

Digital Equipment Corporation, "DEChub 900 MultiSwitch Performance, Multi-Technology Switching Platform Whitepaper," <http://www.networks.digiW.com/html/white-papers.html>. September 1996.

Digital Equipment Corporation, "Digi Networks Frequently Asked Questions on ATM and Digital's ATM program," Whitepaper, <http://www.networks.digital.com/white-papers.html>, June 1996.

Digital Equipment Corporation, "IP Packet Switching on the GIGAswitch/FDDI System Whitepaper,"
<http://www.networks.digital.com/white-papers.html>, January 1997.

Digital Equipment Corporation, "Network Switching: Technology, Strategy and Products" Whitepaper, <http://www.networks.digital.com/white-papers.html>, 1995.

Digital Equipment Corporation, "The Vswitch 900 Family" Whitepaper, <http://www.networks.digital.com/html/white-papers.html>, August 1996.

Duffy, Jim, "New Offerings Better Manage Servers, Frame Relay Nets," Network World, 9 June 1997.

Fisher, Jill E., Establishing a Chargeback Policy: The Department of Defense Can Learn From One Company's Approach, Master's Thesis, Naval Postgraduate School, Monterey, California, December 1993.

Flood, J.H., and others, Telecommunication Networks, Peter Perenrinus Ltd., 1977.

Frame Relay Forum, "Frame Relay: Networks for Tomorrow and Today," URL: <http://www.frforum.com/4000/4001.html>, 1994.

Gaiae, Beth, "The ATM Series: Managing Traffic Flow. 'Network WorldFusion' August 1997.

Gaflant, John, "Chargeback's High Price," Network WorldFusion, v.11, n.50, 12 December 1994.

Graziano, Claudia, "Agents Get Smart LAN Tools Deliver on the Promise of Automated Management," LANTimes Online, 27 February, 1998, URL: <http://www.lantimes.com/lantimes/archive/html>, (23 February 1997).

Grossman, Daniel B., "An Overview of Frame Relay Technology," IEEE, 1991.

Hakulinen, Harold, "Proposed IPv6 Priority Field Semantics," URL: <http://www6.cs-ipv6.lanes.ac.uk/ipv6/mail-archiveAPng/1997-04/0154.html> 14 March 1997.

Henderson, L., and Gage, B., "Stretch Your WAN Limits," Network World Fusion, 11 December 1995.

Herzog, Shai, "RSVP Extensions for Policy Control," URL: <http://ietf.org/internet-drafts/draft-ietf-rsvp-policy-ext-02.txt>, 19 March 1997.

Hewlett-Packard, "LAN-Questions and Answers," URL: <http://hpcc920.external.hp.com/cposupport/networking/support-doc/html>, 1996.

Hibbard, Justin, "IS Looks to Bundled Intranet Sciences," Computerworld, v.31, n.12, 24 March 1997.

Hoffinan, Thomas, "Salomon Brothers Puts Chargeback Online," Computerworld, v.26, n.42, 19 October 1992.

Hudgins-Bonafield, Christine, "Vendor Fall Out over ATM Routing," Network Computing, URL: <http://techweb.cmp.com:80/nc/online/html>

Hurne, Barbara, "Order from Chaos," LANTimes Online, URL: <http://www.lantiimes.com/archive/503bO84a.html>, 27 March 1995.

Ipsilon Products, "IP Switching Applications," URL: <http://www.ipson.com/vproducts/applications.htm>, 31 July 1997.

Johnson, V., Johnson M., and Hall, M., "IP Multicast: Making It Happen," Data Communications, 21 May 1997.

Juliano, Mark, "State of the Art: ATM Traffic Control," Byte, December 1994.

Katzela, and Naghshineh, "Channel Assignment Schemes for Cellular Mobile Tele-communications Systems: A Comprehensive Survey," IETF Personal Communications, June 1996.

Kobielus, James, "Overcoming Net Managers' Fear of Chargeback Systems," Network World Fusion, v.25, n. 1, 17 February 1992.

Kosiur, Dave, "Establish Your Own Management Policy," PC Week, 31 March 1997.

LAN Times Online, "Virtual Concept," URL: <http://www.wcmh.com/lantimes/95dec/512bO22.html>, 8 December 1995,

Laudon, Kenneth C., and Laudon, Lane P., Essentials of Management Information Systems: Organization and Technology, Prentice-Hall, Inc., 1995.

Lisle, Reggie, "Comparison: Software-Metering Tools", LAN Times Online, <http://www.lantimes.com/archive/507aO72b.html>, (23 Feb. 97).

Mackie-Mason, Jeffrey K., "Can Bandwidth Be Reserved?," URL: <http://www.spp.unch.edu/spp/papers/jmm/RAQs/node32.html>, 6 July 1995.

Mackie-Mason, Jeffrey K., and Varian Hal R., "Pricing Congestible Network Resources," <http://www.sims.berkeley.edu/resources/infoecon/Pricing.html>, 11 November 1994.

Mackie-Mason, Jeffery K., and Varian Hal R., "Some Economics of the Internet," [http://www.spp.urnich.edu/spp/paper/Economics-of Intenet.pdf](http://www.spp.urnich.edu/spp/paper/Economics-of%20Intenet.pdf), 17 February 1994.

Madge Networks, "Etheret Switching: A Technology" White Paper, San Jose, CA, 1996.

Madge Networks, "Network Performance and the Client Connection: A Technology," White Paper, May 1997.

Madge Networks, "Solutions Guide to Building a Better Network," Whitepaper, August 1996.

Mclean, Michelle R., "Desktop ATM: Some Power Workgroups Can't Get Enough Bandwidth," LANTimes Online, 16 September 1998.

Mclean, Michelle R., "The Future in Now for Multicast Solution," LANTimes Online, February 1998.

Mclean, Michelle R., "Faster Speeds and Betters Service," LANTimes Online, January 1998.

Mclean, Michelle R., "High-Speed Nightmare," LANTimes Online, March 1998.

Mclean, Michelle R., "Protocol Hype Continues: ATM Users Can Relax-The Differences Between NEOA and I-PNNI are Minor," LANTimes Online, March 1996.

McLean, Mchelle R., "RSVP: Promises and Problems, Limitations Must Be Realized Before Anticipating Real-Time Benefits," LANTimes Online, 2 October 1996.

Muller, N., and Davidson, R., LANs to WANS: Network Management in the 1990s, Artech House, 1990.

Murphy, John, and Murphy, Lian "Bandwidth Allocation By Pricing in ATM Networks," URL: <http://www.eeng.deu.it/~murphy/band-price/band-price.htm>, 8 June 1995.

Murphy, Liam, Murphy, John, and Mackie-Mason, Jeffrey, "Feedback and Efficiency in ATM Networks," URL: <http://www.spp.urffich.edu/spp/papers/jmnVicc96.pdf>, 1996.

Myhrvold, Nathan, "A Penny for Your Thoughts? Charging a Little on the Internet is Even Harder Than Charging a Lot," Slate, URL: <http://www.slate.com/CriticalMass/97-02-13/CriticalMass.asp>, 13 February 1997.

Nagle, John, "Congestion Control in IP/TCP Internetworks," Network Working Group Request For Comments 896, 6 January 1984.

Network General, "How to Manage Switched LANs and ATM Switches for Maximum Performance: A Network Visibility Guide," URL: http://www.ngc.com/white_papers/pdf/SWITCHI.H.PDF, (March 1997).

Network General, "How to Optimize Network Performance While Avoiding Unnecessary investments: A Network Visibility Guide," URL: http://www.ngc.com/white_papers/Optimize/24152.htm, (March 1997).

Network General, "Proactive Solutions to the Five Most Critical Network Problems: A Network Visibility Guide," URL: <http://www.ngc.com/white-Papers/Top5/24158.html>, (March 1997).

Newbold, Paul, Statistics for Business & Economics, Fourth Edition, Prentice Hall, 1995.

Novell, "IP Switching, February 1997, "URL: <http://www.novell.com/nwc/feb.97/switch27.html>.

Ouellette, Tim, "Horizons Unleashes Metering with a Twist," Computerworld, v.29, n.23, (5 June 1995).

Packeteer, "The Technology, March 1997). "URL: <http://www.packeteer.com/technology.htm>, (14 June 1998).

Pancha, P., El Zarki, M., "Prioritized Transmission of Variable Bit Rate NTEG -Video," IEEE, NY, NY, 1992.

Petr, D., Evans, J. Neir, L., Singh, J., and Fronst, V., "Access Traffic Control Implementations for Frame Relay," IEEE International Conference on Communications '93, IEEE, May 1993.

Petrosky, Mary, "Get on Board the Directory Train," 1997. "Network World Fusion, 30 July 1997.

Petrosky, Mary, "Policies: Coming to a Net Near You," Network World Fusion, 7 April 1997.

Pappalardo, Denise, "Frame Relay Gets A New Set of Priorities," Network World Fusion, 20 January 1997.

Platt, A., and Morse, M. J., Some Aspects of Traffic Management in Frame Relay Networks," IEEE Eighth UK Teletraffic Symposium, IEEE, 1991.

Potter, William A., An Analysis of the Navy Regional Automation Center (NARDCD) Chargeback System, Master's Thesis, Naval Postgraduate School, Monterey, California, September 1986.

Raynovich, R. Scott, "Proliferation of Net App Choke Pipes," LANTimes Online, June 1996.

"ReSerVation Protocol (RSVP) Gigabit Network Communication Research (GIGABIT)," <http://www.ito.darpa.mil/Summaries95/8420-USC-ISI-ReSerVation.html>, (23 February 1997).

Roberts, Erica, "Gigabit Ethernet: Weighted Down by Doubts--Is the Proposed High-Speed Spec Too Good to Be True?," Data Communications, November 1996.

Roberts, Erica, "VAN-Fare for the Common LAN," LAN Times Online, January 9, 1995, <http://www.lantimes.com/archive/50laOOla.htm>, (23 March 1997).

Rosenbach, B., and Soref, I., "RMON: the Enterprise Management Standard," Data Communications, URL: <http://www.data.com/Tutorials/Management-Standard.html>, 21 March 1996.

Salamone, S., "Net Traffic Raises Stakes for NCS", LAN Times Online, URL: <http://www.lantimes.com/lantimes/97/97jan/701a030a.html>, January 1997.

Schrage, Michael, "You Get What You Don't Pay For," Computerworld, v.27, n.44, 1 November 1993.

Simmons, Wayne, "Rethinking Systems Management-Consider People, Processes, and Technology When Implementing and Enterprisewide Solution," Information Week, 10 March 1997.

Snell, Monica, "The Price You Pay, Is the Meter Running on Your Intranet? With New Tools, You'll Finally Know," LANTimes Online, March 1997, URL: <http://www.lantimes.com/97/97mar/703b035a.html>, (05 March 1997).

Stedman, Craig, "Where Do You Send the Biff?," Computerworld, v.30, n. 1, 26 December 1995.

Stem, Dan, and Mazelia, Frank, "...On the Subject of ATM," URL: <http://www.datacomm-us.com/technow/scan06/scan06.html>, 1997.

Stem, Dan, and Mazelia, Frank, "...On the Subject of Routers", URL: <http://www.datacomm-us.com/technow/scan01/scan01.html>, 1997.

Taylor, Martin, "LAN Emulation Over ATM: A Technology," White Paper, November 1998.

Tekinay, S., and Jabbari, B., "Analysis of Measurement Based Prioritization Schemes for Handovers in Cellular Networks," IEEE, NY, NY., 1992.

The Economist, "HangOn," URL: <http://www.economist.com/issue/19-10-96/sf0775.html>, (19 Oct 1998).

The Economist, "The Economics of the Internet: Too Cheap to Meter?," URL: <http://www.economist.com/issue/19-10-96/sf0774.html>, (19 Oct 1996).

The Economist, "Why the Net Should Grow Up," URL: <http://www.economist.com/issue/19-10-96/ld4401.html>, (19 Oct 1996).

Trovini, Kevin L., Analysis of Network Traffic and Bandwidth Capacity: LoadBalancing and Rightsizing of Wide Area Network Links, Master's Thesis, Naval Postgraduate School, September 1996.

UNH InterOperability Lab, "Demand Priority Protocol," URL: <http://www.iol.unh.edu/training/vganylan/mac/demandpr.html>, 1997.

Van Norman, Harrell, LAN-WAN Optimization Techniques, Artech House, 1997.

Wace, Bob, "User Response Weak on RSVP," Computerworld, 31 March 1997.

Wayner, Peter, "Time and Money," Byte, April 1990.

Wirbel, Loring, "Asynchronous Transfer Mode Threatened at Own Confab-ATM Switching Takes It on the Chin," Electrical Engineering Times, 13 May 1996.

Wirbel, Loring, "Cajun ASICs Spice Gbit Ethernet Nfix," Electrical Engineering Times, 3 February 1997.

INITIAL DISTRIBUTION LIST

1. Defense Technical Information Center2
8725 John J. Kingman Road, Ste 0944
Fort Belvoir, VA 22060-6218

2. Dudley Knox Library2
Naval Postgraduate School
411 Dyer Road
Monterey, California 93943-5101

3. Professor Geoff Xie, Code CS2
Naval Postgraduate School
Monterey, CA 93943

4. Professor Bert Lundy, Code CS2
Naval Postgraduate School
Monterey, CA 93943

5. LCDR Paul Wetzel.....3
514 Hindsdale Dr.
Arlington, TX 76006

6. Commanding Officer, NCTS Sicily.....2
PSC 812 Box 3290
FPO AE 09427

7. Computational Sciences Division.1
NASA Ames Research Center
Attn: Marjori Johnson, Senior Scientist
MS 269-2 Moffett Field, CA 94035-1000

8. Chairman, Code IS.1
Naval Postgraduate School
Monterey, CA 93943