

Network Working Group  
Request for Comments: 1287

D. Clark  
MIT  
L. Chapin  
BBN  
V. Cerf  
CNRI  
R. Braden  
ISI  
R. Hobby  
UC Davis  
December 1991

## Towards the Future Internet Architecture

### Status of this Memo

This informational RFC discusses important directions for possible future evolution of the Internet architecture, and suggests steps towards the desired goals. It is offered to the Internet community for discussion and comment. This memo provides information for the Internet community. It does not specify an Internet standard. Distribution of this memo is unlimited.

### Table of Contents

|                                       |    |
|---------------------------------------|----|
| 1. INTRODUCTION .....                 | 2  |
| 2. ROUTING AND ADDRESSING .....       | 5  |
| 3. MULTI-PROTOCOL ARCHITECTURES ..... | 9  |
| 4. SECURITY ARCHITECTURE .....        | 13 |
| 5. TRAFFIC CONTROL AND STATE .....    | 16 |
| 6. ADVANCED APPLICATIONS .....        | 18 |
| 7. REFERENCES .....                   | 21 |
| APPENDIX A. Setting the Stage .....   | 22 |
| APPENDIX B. Group Membership .....    | 28 |
| Security Considerations .....         | 29 |
| Authors' Addresses .....              | 29 |

## 1. INTRODUCTION

### 1.1 The Internet Architecture

The Internet architecture, the grand plan behind the TCP/IP protocol suite, was developed and tested in the late 1970s by a small group of network researchers [1-4]. Several important features were added to the architecture during the early 1980's -- subnetting, autonomous systems, and the domain name system [5,6]. More recently, IP multicasting has been added [7].

Within this architectural framework, the Internet Engineering Task Force (IETF) has been working with great energy and effectiveness to engineer, define, extend, test, and standardize protocols for the Internet. Three areas of particular importance have been routing protocols, TCP performance, and network management. Meanwhile, the Internet infrastructure has continued to grow at an astonishing rate. Since January 1983 when the ARPANET first switched from NCP to TCP/IP, the vendors, managers, wizards, and researchers of the Internet have all been laboring mightily to survive their success.

A set of the researchers who had defined the Internet architecture formed the original membership of the Internet Activities Board (IAB). The IAB evolved from a technical advisory group set up in 1981 by DARPA to become the general technical and policy oversight body for the Internet. IAB membership has changed over the years to better represent the changing needs and issues in the Internet community, and more recently, to reflect the internationalization of the Internet, but it has retained an institutional concern for the protocol architecture.

The IAB created the Internet Engineering Task Force (IETF) to carry out protocol development and engineering for the Internet. To manage the burgeoning IETF activities, the IETF chair set up the Internet Engineering Steering Group (IESG) within the IETF. The IAB and IESG work closely together in ratifying protocol standards developed within the IETF.

Over the past few years, there have been increasing signs of strains on the fundamental architecture, mostly stemming from continued Internet growth. Discussions of these problems reverberate constantly on many of the major mailing lists.

### 1.2 Assumptions

The priority for solving the problems with the current Internet architecture depends upon one's view of the future relevance of

TCP/IP with respect to the OSI protocol suite. One view has been that we should just let the TCP/IP suite strangle in its success, and switch to OSI protocols. However, many of those who have worked hard and successfully on Internet protocols, products, and service are anxious to try to solve the new problems within the existing framework. Furthermore, some believe that OSI protocols will suffer from versions of many of the same problems.

To begin to attack these issues, the IAB and the IESG held a one-day joint discussion of Internet architectural issues in January 1991. The framework for this meeting was set by Dave Clark (see Appendix A for his slides). The discussion was spirited, provocative, and at times controversial, with a lot of soul-searching over questions of relevance and future direction. The major result was to reach a consensus on the following four basic assumptions regarding the networking world of the next 5-10 years.

- (1) The TCP/IP and OSI suites will coexist for a long time.

There are powerful political and market forces as well as some technical advantages behind the introduction of the OSI suite. However, the entrenched market position of the TCP/IP protocols means they are very likely to continue in service for the foreseeable future.

- (2) The Internet will continue to include diverse networks and services, and will never be comprised of a single network technology.

Indeed, the range of network technologies and characteristics that are connected into the Internet will increase over the next decade.

- (3) Commercial and private networks will be incorporated, but we cannot expect the common carriers to provide the entire service. There will be mix of public and private networks, common carriers and private lines.

- (4) The Internet architecture needs to be able to scale to  $10^9$  networks.

The historic exponential growth in the size of the Internet will presumably saturate some time in the future, but forecasting when is about as easy as forecasting the future economy. In any case, responsible engineering requires an architecture that is CAPABLE of expanding to a worst-case size. The exponent "9" is rather fuzzy; estimates have varied from 7 to 10.

### 1.3 Beginning a Planning Process

Another result of the IAB and IESG meeting was the following list of the five most important areas for architectural evolution:

(1) Routing and Addressing

This is the most urgent architectural problem, as it is directly involved in the ability of the Internet to continue to grow successfully.

(2) Multi-Protocol Architecture

The Internet is moving towards widespread support of both the TCP/IP and the OSI protocol suites. Supporting both suites raises difficult technical issues, and a plan -- i.e., an architecture -- is required to increase the chances of success. This area was facetiously dubbed "making the problem harder for the good of mankind."

Clark had observed that translation gateways (e.g., mail gateways) are very much a fact of life in Internet operation but are not part of the architecture or planning. The group discussed the possibility of building the architecture around the partial connectivity that such gateways imply.

(3) Security Architecture

Although military security was considered when the Internet architecture was designed, the modern security issues are much broader, encompassing commercial requirements as well. Furthermore, experience has shown that it is difficult to add security to a protocol suite unless it is built into the architecture from the beginning.

(4) Traffic Control and State

The Internet should be extended to support "real-time" applications like voice and video. This will require new packet queueing mechanisms in gateways -- "traffic control" -- and additional gateway state.

(5) Advanced Applications

As the underlying Internet communication mechanism matures, there is an increasing need for innovation and standardization in building new kinds of applications.

The IAB and IESG met again in June 1991 at SDSC and devoted three full days to a discussion of these five topics. This meeting, which was called somewhat perversely the "Architecture Retreat", was convened with a strong resolve to take initial steps towards planning evolution of the architecture. Besides the IAB and IESG, the group of 32 people included the members of the Research Steering Group (IRSG) and a few special guests. On the second day, the Retreat broke into groups, one for each of the five areas. The group membership is listed in Appendix B.

This document was assembled from the reports by the chairs of these groups. This material was presented at the Atlanta IETF meeting, and appears in the minutes of that meeting [8].

## 2. ROUTING AND ADDRESSING

Changes are required in the addressing and routing structure of IP to deal with the anticipated growth and functional evolution of the Internet. We expect that:

- o The Internet will run out of certain classes of IP network addresses, e.g., B addresses.
- o The Internet will run out of the 32-bit IP address space altogether, as the space is currently subdivided and managed.
- o The total number of IP network numbers will grow to the point where reasonable routing algorithms will not be able to perform routing based upon network numbers.
- o There will be a need for more than one route from a source to a destination, to permit variation in TOS and policy conformance. This need will be driven both by new applications and by diverse transit services. The source, or an agent acting for the source, must control the selection of the route options.

### 2.1 Suggested Approach

There is general agreement on the approach needed to deal with these facts.

- (a) We must move to an addressing scheme in which network numbers are aggregated into larger units as the basis for routing. An example of an aggregate is the Autonomous System, or the Administrative Domain (AD).

Aggregation will accomplish several goals: define regions where policy is applied, control the number of routing

elements, and provide elements for network management. Some believe that it must be possible to further combine aggregates, as in a nesting of ADs.

- (b) We must provide some efficient means to compute common routes, and some general means to compute "special" routes.

The general approach to special routes will be some form of route setup specified by a "source route".

There is not full agreement on how ADs may be expected to be aggregated, or how routing protocols should be organized to deal with the aggregation boundaries. A very general scheme may be used [ref. Chiappa], but some prefer a scheme that more restricts and defines the expected network model.

To deal with the address space exhaustion, we must either expand the address space or else reuse the 32 bit field ("32bf") in different parts of the net. There are several possible address formats that might make sense, as described in the next section.

Perhaps more important is the question of how to migrate to the new scheme. All migration plans will require that some routers (or other components inside the Internet) be able to rewrite headers to accommodate hosts that handle only the old or format or only the new format. Unless the need for such format conversion can be inferred algorithmically, migration by itself will require some sort of setup of state in the conversion element.

We should not plan a series of "small" changes to the architecture. We should embark now on a plan that will take us past the exhaustion of the address space. This is a more long-range act of planning than the Internet community has undertaken recently, but the problems of migration will require a long lead time, and it is hard to see an effective way of dealing with some of the more immediate problems, such as class B exhaustion, in a way that does not by itself take a long time. So, once we embark on a plan of change, it should take us all the way to replacing the current 32-bit global address space. (This conclusion is subject to revision if, as is always possible, some very clever idea surfaces that is quick to deploy and gives us some breathing room. We do not mean to discourage creative thinking about short-term actions. We just want to point out that even small changes take a long time to deploy.)

Conversion of the address space by itself is not enough. We must at the same time provide a more scalable routing architecture, and tools to better manage the Internet. The proposed approach is to

ADs as the unit of aggregation for routing. We already have partial means to do this. IDPR does this. The OSI version of BGP (IDRP) does this. BGP could evolve to do this. The additional facility needed is a global table that maps network numbers to ADs.

For several reasons (special routes and address conversion, as well as accounting and resource allocation), we are moving from a "stateless" gateway model, where only precomputed routes are stored in the gateway, to a model where at least some of the gateways have per-connection state.

## 2.2 Extended IP Address Formats

There are three reasonable choices for the extended IP address format.

- A) Replace the 32 bit field (32bf) with a field of the same size but with different meaning. Instead of being globally unique, it would now be unique only within some smaller region (an AD or an aggregate of ADs). Gateways on the boundary would rewrite the address as the packet crossed the boundary.

Issues: (1) addresses in the body of packets must be found and rewritten; (2) the host software need not be changed; (3) some method (perhaps a hack to the DNS) must set up the address mappings.

This scheme is due to Van Jacobson. See also the work by Paul Tsuchiya on NAT.

- B) Expand the 32bf to a 64 bit field (or some other new size), and use the field to hold a global host address and an AD for that host.

This choice would provide a trivial mapping from the host to the value (the AD) that is the basis of routing. Common routes (those selected on the basis of destination address without taking into account the source address as well) can be selected directly from the packet address, as is done today, without any prior setup.

- 3) Expand the 32bf to a 64 bit field (or some other new size), and use the field as a "flat" host identifier. Use connection setup to provide routers with the mapping from host id to AD, as needed.

The 64 bits can now be used to simplify the problem of allocating host ids, as in Ethernet addresses.

Each of these choices would require an address re-writing module as a part of migration. The second and third require a change to the IP header, so host software must change.

### 2.3 Proposed Actions

The following actions are proposed:

#### A) Time Line

Construct a specific set of estimates for the time at which the various problems above will arise, and construct a corresponding time-line for development and deployment of a new addressing/routing architecture. Use this time line as a basis for evaluating specific proposals for changes. This is a matter for the IETF.

#### B) New Address Format

Explore the options for a next generation address format and develop a plan for migration. Specifically, construct a prototype gateway that does address mapping. Understand the complexity of this task, to guide our thinking about migration options.

#### C) Routing on ADs

Take steps to make network aggregates (ADs) the basis of routing. In particular, explore the several options for a global table that maps network numbers to ADs. This is a matter for the IETF.

#### D) Policy-Based Routing

Continue the current work on policy based routing. There are several specific objectives.

- Seek ways to control the complexity of setting policy (this is a human interface issue, not an algorithm complexity issue).
- Understand better the issues of maintaining connection state in gateways.
- Understand better the issues of connection state setup.

### E) Research on Further Aggregation

Explore, as a research activity, how ADs should be aggregated into still larger routing elements.

- Consider whether the architecture should define the "role" of an AD or an aggregate.
- Consider whether one universal routing method or distinct methods should be used inside and outside ADs and aggregates.

Existing projects planned for DARTnet will help resolve several of these issues: state in gateways, state setup, address mapping, accounting and so on. Other experiments in the R&D community also bear on this area.

## 3. MULTI-PROTOCOL ARCHITECTURE

Changing the Internet to support multiple protocol suites leads to three specific architectural questions:

- o How exactly will we define "the Internet"?
- o How would we architect an Internet with  $n > 1$  protocol suites, regardless of what the suites are?
- o Should we architect for partial or filtered connectivity?
- o How to add explicit support for application gateways into the architecture?

### 3.1 What is the "Internet"?

It is very difficult to deal constructively with the issue of "the multi-protocol Internet" without first determining what we believe "the Internet" is (or should be). We distinguish "the Internet", a set of communicating systems, from "the Internet community", a set of people and organizations. Most people would accept a loose definition of the latter as "the set of people who believe themselves to be part of the Internet community". However, no such "sociological" definition of the Internet itself is likely to be useful.

Not too long ago, the Internet was defined by IP connectivity (IP and ICMP were - and still are - the only "required" Internet protocols). If I could PING you, and you could PING me, then we were both on the Internet, and a satisfying working definition of

the Internet could be constructed as a roughly transitive closure of IP-speaking systems. This model of the Internet was simple, uniform, and - perhaps most important - testable. The IP-connectivity model clearly distinguished systems that were "on the Internet" from those that were not.

As the Internet has grown and the technology on which it is based has gained widespread commercial acceptance, the sense of what it means for a system to be "on the Internet" has changed, to include:

- \* Any system that has partial IP connectivity, restricted by policy filters.
- \* Any system that runs the TCP/IP protocol suite, whether or not it is actually accessible from other parts of the Internet.
- \* Any system that can exchange RFC-822 mail, without the intervention of mail gateways or the transformation of mail objects.
- \* Any system with e-mail connectivity to the Internet, whether or not a mail gateway or mail object transformation is required.

These definitions of "the Internet", are still based on the original concept of connectivity, just "moving up the stack".

We propose instead a new definition of the Internet, based on a different unifying concept:

- \* "Old" Internet concept: IP-based.  
  
The organizing principle is the IP address, i.e., a common network address space.
- \* "New" Internet concept: Application-based.  
  
The organizing principle is the domain name system and directories, i.e., a common - albeit necessarily multiform - application namespace.

This suggests that the idea of "connected status", which has traditionally been tied to the IP address(via network numbers, should instead be coupled to the names and related identifying information contained in the distributed Internet directory.

A naming-based definition of "the Internet" implies a much larger Internet community, and a much more dynamic (and unpredictable) operational Internet. This argues for an Internet architecture based on adaptability (to a broad spectrum of possible future developments) rather than anticipation.

### 3.2 A Process-Based Model of the Multiprotocol Internet

Rather than specify a particular "multi-protocol Internet", embracing a pre-determined number of specific protocol architectures, we propose instead a process-oriented model of the Internet, which accommodates different protocol architectures according to the traditional "things that work" principle.

A process-oriented Internet model includes, as a basic postulate, the assertion that there is no \*steady-state\* "multi-protocol Internet". The most basic forces driving the evolution of the Internet are pushing it not toward multi-protocol diversity, but toward the original state of protocol-stack uniformity (although it is unlikely that it will ever actually get there). We may represent this tendency of the Internet to evolve towards homogeneity as the most "thermodynamically stable" state by describing four components of a new process-based Internet architecture:

#### Part 1: The core Internet architecture

This is the traditional TCP/IP-based architecture. It is the "magnetic center" of Internet evolution, recognizing that (a) homogeneity is still the best way to deal with diversity in an internetwork, and (b) IP connectivity is still the best basic model of the Internet (whether or not the actual state of IP ubiquity can be achieved in practice in a global operational Internet).

"In the beginning", the Internet architecture consisted only of this first part. The success of the Internet, however, has carried it beyond its uniform origins; ubiquity and uniformity have been sacrificed in order to greatly enrich the Internet "gene pool".

Two additional parts of the new Internet architecture express the ways in which the scope and extent of the Internet have been expanded.

#### Part 2: Link sharing

Here physical resources -- transmission media, network

interfaces, perhaps some low-level (link) protocols -- are shared by multiple, non-interacting protocol suites. This part of the architecture recognizes the necessity and convenience of coexistence, but is not concerned with interoperability; it has been called "ships in the night" or "S.I.N.".

Coexisting protocol suites are not, of course, genuinely isolated in practice; the ships passing in the night raise issues of management, non-interference, coordination, and fairness in real Internet systems.

### Part 3: Application interoperability

Absent ubiquity of interconnection (i.e., interoperability of the "underlying stacks"), it is still possible to achieve ubiquitous application functionality by arranging for the essential semantics of applications to be conveyed among disjoint communities of Internet systems. This can be accomplished by application relays, or by user agents that present a uniform virtual access method to different application services by expressing only the shared semantics.

This part of the architecture emphasizes the ultimate role of the Internet as a basis for communication among applications, rather than as an end in itself. To the extent that it enables a population of applications and their users to move from one underlying protocol suite to another without unacceptable loss of functionality, it is also a "transition enabler".

Adding parts 2 and 3 to the original Internet architecture is at best a mixed blessing. Although they greatly increase the scope of the Internet and the size of the Internet community, they also introduce significant problems of complexity, cost, and management, and they usually represent a loss of functionality (particularly with respect to part 3). Parts 2 and 3 represent unavoidable, but essentially undesirable, departures from the homogeneity represented by part 1. Some functionality is lost, and additional system complexity and cost is endured, in order to expand the scope of the Internet. In a perfect world, however, the Internet would evolve and expand without these penalties.

There is a tendency, therefore, for the Internet to evolve in favor of the homogeneous architecture represented by part 1, and away from the compromised architectures of parts 2 and 3. Part 4 expresses this tendency.

#### Part 4: Hybridization/Integration.

Part 4 recognizes the desirability of integrating similar elements from different Internet protocol architectures to form hybrids that reduce the variability and complexity of the Internet system. It also recognizes the desirability of leveraging the existing Internet infrastructure to facilitate the absorption of "new stuff" into the Internet, applying to "new stuff" the established Internet practice of test, evaluate, adopt.

This part expresses the tendency of the Internet, as a system, to attempt to return to the original "state of grace" represented by the uniform architecture of part 1. It is a force acting on the evolution of the Internet, although the Internet will never actually return to a uniform state at any point in the future.

According to this dynamic process model, running X.400 mail over RFC 1006 on a TCP/IP stack, integrated IS-IS routing, transport gateways, and the development of a single common successor to the IP and CLNP protocols are all examples of "good things". They represent movement away from the non-uniformity of parts 2 and 3 towards greater homogeneity, under the influence of the "magnetic field" asserted by part 1, following the hybridization dynamic of part 4.

## 4. SECURITY ARCHITECTURE

### 4.1 Philosophical Guidelines

The principal themes for development of an Internet security architecture are simplicity, testability, trust, technology and security perimeter identification.

- \* There is more to security than protocols and cryptographic methods.
- \* The security architecture and policies should be simple enough to be readily understood. Complexity breeds misunderstanding and poor implementation.
- \* The implementations should be testable to determine if the policies are met.
- \* We are forced to trust hardware, software and people to make any security architecture function. We assume that the technical instruments of security policy enforcement are at

least as powerful as modern personal computers and work stations; we do not require less capable components to be self-protecting (but might apply external remedies such as link level encryption devices).

- \* Finally, it is essential to identify security perimeters at which protection is to be effective.

#### 4.2 Security Perimeters

There were four possible security perimeters: link level, net/subnet level, host level, and process/application level. Each imposes different requirements, can admit different techniques, and makes different assumptions about what components of the system must be trusted to be effective.

Privacy Enhanced Mail is an example of a process level security system; providing authentication and confidentiality for SNMP is another example. Host level security typically means applying an external security mechanism on the communication ports of a host computer. Network or subnetwork security means applying the external security capability at the gateway/router(s) leading from the subnetwork to the "outside". Link-level security is the traditional point-to-point or media-level (e.g., Ethernet) encryption mechanism.

There are many open questions about network/subnetwork security protection, not the least of which is a potential mismatch between host level (end/end) security methods and methods at the network/subnetwork level. Moreover, network level protection does not deal with threats arising within the security perimeter.

Applying protection at the process level assumes that the underlying scheduling and operating system mechanisms can be trusted not to prevent the application from applying security when appropriate. As the security perimeter moves downward in the system architecture towards the link level, one must make many assumptions about the security threat to make an argument that enforcement at a particular perimeter is effective. For example, if only link-level encryption is used, one must assume that attacks come only from the outside via communications lines, that hosts, switches and gateways are physically protected, and the people and software in all these components are to be trusted.

#### 4.3 Desired Security Services

We need authenticatable distinguished names if we are to implement discretionary and non-discretionary access control at application

and lower levels in the system. In addition, we need enforcement for integrity (anti-modification, anti-spoof and anti-replay defenses), confidentiality, and prevention of denial-of-service. For some situations, we may also need to prevent repudiation of message transmission or to prevent covert channels.

We have some building blocks with which to build the Internet security system. Cryptographic algorithms are available (e.g., Data Encryption Standard, RSA, El Gamal, and possibly other public key and symmetric key algorithms), as are hash functions such as MD2 and MD5.

We need Distinguished Names (in the OSI sense) and are very much in need of an infrastructure for the assignment of such identifiers, together with widespread directory services for making them known. Certificate concepts binding distinguished names to public keys and binding distinguished names to capabilities and permissions may be applied to good advantage.

At the router/gateway level, we can apply address and protocol filters and other configuration controls to help fashion a security system. The proposed OSI Security Protocol 3 (SP3) and Security Protocol 4 (SP4) should be given serious consideration as possible elements of an Internet security architecture.

Finally, it must be observed that we have no good solutions to safely storing secret information (such as the secret component of a public key pair) on systems like PCs or laptop computers that are not designed to enforce secure storage.

#### 4.4 Proposed Actions

The following actions are proposed.

##### A) Security Reference Model

A Security Reference Model for the Internet is needed, and it should be developed expeditiously. This model should establish the target perimeters and document the objectives of the security architecture.

##### B) Privacy-Enhanced Mail (PEM)

For Privacy Enhanced Mail, the most critical steps seem to be the installation of (1) a certificate generation and management infrastructure, and (2) X.500 directory services to provide access to public keys via distinguished names. Serious attention also needs to be placed on any limitations

imposed by patent and export restrictions on the deployment of this system.

C) Distributed System Security

We should examine security methods for distributed systems applications, in both simple (client/server) and complex (distributed computing environment) cases. For example, the utility of certificates granting permissions/capabilities to objects bound to distinguished names should be examined.

D) Host-Level Security

SP4 should be evaluated for host-oriented security, but SP3 should also be considered for this purpose.

E) Application-Level Security

We should implement application-level security services, both for their immediate utility (e.g., PEM, SNMP authentication) and also to gain valuable practical experience that can inform the refinement of the Internet security architecture.

## 5. TRAFFIC CONTROL AND STATE

In the present Internet, all IP datagrams are treated equally. Each datagram is forwarded independently, regardless of any relationship it has to other packets for the same connection, for the same application, for the same class of applications, or for the same user class. Although Type-of-Service and Precedence bits are defined in the IP header, these are not generally implemented, and in fact it is not clear how to implement them.

It is now widely accepted that the future Internet will need to support important applications for which best-effort is not sufficient -- e.g., packet video and voice for teleconferencing. This will require some "traffic control" mechanism in routers, controlled by additional state, to handle "real-time" traffic.

### 5.1 Assumptions and Principles

- o ASSUMPTION: The Internet will need to support performance guarantees for particular subsets of the traffic.

Unfortunately, we are far from being able to give precise meanings to the terms "performance", "guarantees", or "subsets" in this statement. Research is still needed to answer these questions.

- o The default service will continue to be the current "best-effort" datagram delivery, with no service guarantees.
- o The mechanism of a router can be separated into (1) the forwarding path and (2) the control computations (e.g., routing) which take place in the background.

The forwarding path must be highly optimized, sometimes with hardware-assist, and it is therefore relatively costly and difficult to change. The traffic control mechanism operates in the forwarding path, under the control of state created by routing and resource control computations that take place in background. We will have at most one shot at changing the forwarding paths of routers, so we had better get it right the first time.

- o The new extensions must operate in a highly heterogeneous environment, in which some parts will never support guarantees. For some hops of a path (e.g., a high-speed LAN), "over-provisioning" (i.e., excess capacity) will allow adequate service for real-time traffic, even when explicit resource reservation is unavailable.
- o Multicast distribution is probably essential.

## 5.2 Technical Issues

There are a number of technical issues to be resolved, including:

- o Resource Setup

To support real-time traffic, resources need to be reserved in each router along the path from source to destination. Should this new router state be "hard" (as in connections) or "soft" (i.e., cached state)?

- o Resource binding vs. route binding

Choosing a path from source to destination is traditionally performed using a dynamic routing protocol. The resource binding and the routing might be folded into a single complex process, or they might be performed essentially independently. There is a tradeoff between complexity and efficiency.

- o Alternative multicast models

IP multicasting uses a model of logical addressing in which

targets attach themselves to a group. In ST-2, each host in a multicast session includes in its setup packet an explicit list of target addresses. Each of these approaches has advantages and drawbacks; it is not currently clear which will prevail for n-way teleconferences.

- o Resource Setup vs. Inter-AD routing

Resource guarantees of whatever flavor must hold across an arbitrary end-to-end path, including multiple ADs. Hence, any resource setup mechanism needs to mesh smoothly with the path setup mechanism incorporated into IDPR.

- o Accounting

The resource guarantee subsets ("classes") may be natural units for accounting.

### 5.3 Proposed Actions

The actions called for here are further research on the technical issues listed above, followed by development and standardization of appropriate protocols. DARTnet, the DARPA Research Testbed network, will play an important role in this research.

## 6. ADVANCED APPLICATIONS

One may ask: "What network-based applications do we want, and why don't we have them now?" It is easy to develop a large list of potential applications, many of which would be based on a client/server model. However, the more interesting part of the question is: "Why haven't people done them already?" We believe the answer to be that the tools to make application writing easy just do not exist.

To begin, we need a set of common interchange formats for a number of data items that will be used across the network. Once these common data formats have been defined, we need to develop tools that the applications can use to move the data easily.

### 6.1 Common Interchange Formats

The applications have to know the format of information that they are exchanging, for the information to have any meaning. The following format types are to concern:

- (1) Text - Of the formats in this list, text is the most stable, but today's international Internet has to address the needs

of character sets other than USASCII.

- (2) Image - As we enter the "Multimedia Age", images will become increasingly important, but we need to agree on how to represent them in packets.
- (3) Graphics - Like images, vector graphic information needs a common definition. With such a format we could exchange things like architectural blueprints.
- (4) Video - Before we can have a video window running on our workstation, we need to know the format of that video information coming over the network.
- (5) Audio/Analog - Of course, we also need the audio to go with the video, but such a format would be used for representation of all types of analog signals.
- (6) Display - Now that we are opening windows on our workstation, we want to open a window on another person's workstation to show her some data pertinent to the research project, so now we need a common window display format.
- (7) Data Objects - For inter-process communications we need to agree on the formats of things like integers, reals, strings, etc.

Many of these formats are being defined by other, often several other, standards organizations. We need to agree on one format per category for the Internet.

## 6.2 Data Exchange Methods

Applications will require the following methods of data exchange.

### (1) Store and Forward

Not everyone is on the network all the time. We need a standard means of providing an information flow to sometimes-connected hosts, i.e., we need a common store-and-forward service. Multicasting should be included in such a service.

### (2) Global File Systems

Much of the data access over the network can be broken down to simple file access. If you had a real global file system where you access any file on the Internet (assuming you have

permission), would you ever need FTP?

(3) Inter-process Communications

For a true distributed computing environment, we need the means to allow processes to exchange data in a standard method over the network. This requirement encompasses RPC, APIs, etc.

(4) Data Broadcast

Many applications need to send the same information to many other hosts. A standard and efficient method is needed to accomplish this.

(5) Database Access

For good information exchange, we need to have a standard means for accessing databases. The Global File System can get you to the data, but the database access methods will tell you about its structure and content.

Many of these items are being addressed by other organizations, but for Internet interoperability, we need to agree on the methods for the Internet.

Finally, advanced applications need solutions to the problems of two earlier areas in this document. From the Traffic Control and State area, applications need the ability to transmit real-time data. This means some sort of expectation level for data delivery within a certain time frame. Applications also require global authentication and access control systems from the Security area. Much of the usefulness of today's Internet applications is lost due to the lack of trust and security. This needs to be solved for tomorrow's applications.

## 7. REFERENCES

- [1] Cerf, V. and R. Kahn, "A Protocol for Packet Network Intercommunication," IEEE Transactions on Communication, May 1974.
- [2] Postel, J., Sunshine, C., and D. Cohen, "The ARPA Internet Protocol," Computer Networks, Vol. 5, No. 4, July 1981.
- [3] Leiner, B., Postel, J., Cole, R., and D. Mills, "The DARPA Internet Protocol Suite," Proceedings INFOCOM 85, IEEE, Washington DC, March 1985. Also in: IEEE Communications Magazine, March 1985.
- [4] Clark, D., "The Design Philosophy of the DARPA Internet Protocols", Proceedings ACM SIGCOMM '88, Stanford, California, August 1988.
- [5] Mogul, J., and J. Postel, "Internet Standard Subnetting Procedure", RFC 950, USC/Information Sciences Institute, August 1985.
- [6] Mockapetris, P., "Domain Names - Concepts and Facilities", RFC 1034, USC/Information Sciences Institute, November 1987.
- [7] Deering, S., "Host Extensions for IP Multicasting", RFC 1112, Stanford University, August 1989.
- [8] "Proceedings of the Twenty-First Internet Engineering Task Force", Bell-South, Atlanta, July 29 - August 2, 1991.

APPENDIX A: Setting the Stage

Slide 1

WHITHER THE INTERNET?  
OPTIONS FOR ARCHITECTURE

IAB/IESG -- Jan 1990

David D. Clark

---

Slide 2

SETTING THE TOPIC OF DISCUSSION

Goals:

- o Establish a common frame of understanding for IAB, IESG and the Internet community.
- o Understand the set of problems to be solved.
- o Understand the range of solutions open to us.
- o Draw some conclusions, or else "meta-conclusions".

---

Slide 3

## SOME CLAIMS -- MY POSITION

We have two different goals:

- o Make it possible to build "The Internet"
- o Define a protocol suite called Internet

Claim: These goals have very different implications.

The protocols are but a means, though a powerful one.

Claim: If "The Internet" is to succeed and grow, it will require specific design efforts. This need will continue for at least another 10 years.

Claim: Uncontrolled growth could lead to chaos.

Claim: A grass-roots solution seems to be the only means to success. Top-down mandates are powerless.

---

Slide 4

## OUTLINE OF PRESENTATION

- 1) The problem space and the solution space.
- 2) A set of specific questions -- discussion.
- 3) Return to top-level questions -- discussion.
- 4) Plan for action -- meta discussion.

Try to separate functional requirements from technical approach.

Understand how we are bounded by our problem space and our solution space.

Is architecture anything but protocols?

---

Slide 5

## WHAT IS THE PROBLEM SPACE?

Routing and addressing:

How big, what topology, and what routing model?

Getting big:

User services, what technology for host and nets?

Divestiture of the Internet:

Accounting, controlling usage and fixing faults.

New services:

Video? Transactions? Distributed computing?

Security:

End node or network? Routers or relays?

---

Slide 6

## BOUNDING THE SOLUTION SPACE

How far can we migrate from the current state?

- o Can we change the IP header (except to OSI)?
- o Can we change host requirements in mandatory ways?
- o Can we manage a long-term migration objective?
  - Consistent direction vs. diverse goals, funding.

Can we assume network-level connectivity?

- o Relays are the wave of the future (?)
- o Security a key issue; along with conversion.
- o Do we need a new "relay-based" architecture?

How "managed" can/must "The Internet" be?

- o Can we manage or constrain connectivity?

What protocols are we working with? One or many?

---

Slide 7

## THE MULTI-PROTOCOL INTERNET

"Making the problem harder for the good of mankind."

Are we migrating, interoperating, or tolerating multiple protocols?

- o Not all protocol suites will have same range of functionality at the same time.
- o "The Internet" will require specific functions.

Claim: Fundamental conflict (not religion or spite):

- o Meeting aggressive requirements for the Internet
- o Dealing with OSI migration.

Conclusion: One protocol must "lead", and the others must follow.  
When do we "switch" to OSI?

Consider every following slide in this context.

---

Slide 8

## ROUTING and ADDRESSING

What is the target size of "The Internet"?

- o How do addresses and routes relate?
- o What is the model of topology?
- o What solutions are possible?

What range of policy routing is required?

- o BGP and IDRP are two answers. What is the question?
- o Fixed classes, or variable paths?
- o Source controlled routing is a minimum.

How seamless is the needed support for mobile hosts?

- o New address class, rebind to local address, use DNS?

Shall we push for Internet multicast?

---

Slide 9

## GETTING BIG -- AN OLD TITLE

(Addressing and routing was on previous slide...)

What user services will be needed in the next 10 years?

- o Can we construct a plan?
- o Do we need architectural changes?

Is there a requirement for dealing better with ranges in speed, packet sizes, etc.

- o Policy to phase out fragmentation?

What range of hosts (things != Unix) will we support?

---

Slide 10

## DEALING WITH DIVESTITURE

The Internet is composed of parts separately managed and controlled.

What support is needed for network charging?

- o No architecture implies bulk charges and re-billing, pay for lost packets.
- o Do we need controls to supply billing id or routing?

Requirement: we must support links with controlled sharing.  
(Simple form is classes based on link id.)

- o How general?

Is there an increased need for fault isolation? (I vote yes!)

- o How can we find managers to talk to?
- o Do we need services in hosts?

---

Slide 11

NEW SERVICES

- Shall we support video and audio? Real time? What %?
- o Need to plan for input from research. What quality?
  - o Target date for heads-up to vendors.
- Shall we "better" support transactions?
- o Will TCP do? VMTP? Presentation? Locking?
- What application support veneers are coming?
- o Distributed computing -- will it actually happen?
  - o Information networking?

---

Slide 12

SECURITY

- Can we persist in claiming the end-node is the only line of defense?
- o What can we do inside the network?
  - o What can ask the host to do?
- Do we tolerate relays, or architect them?
- Can find a better way to construct security boundaries?
- Do we need global authentication?
- Do we need new host requirements:
- o Logging.
  - o Authentication.
  - o Management interfaces.
    - Phone number or point of reference.

## APPENDIX B: Group Membership

## Group 1: ROUTING AND ADDRESSING

Dave Clark, MIT [Chair]  
Hans-Werner Braun, SDSC  
Noel Chiappa, Consultant  
Deborah Estrin, USC  
Phill Gross, CNRI  
Bob Hinden, BBN  
Van Jacobson, LBL  
Tony Lauck, DEC.

## Group 2: MULTI-PROTOCOL ARCHITECTURE

Lyman Chapin, BBN [Chair]  
Ross Callon, DEC  
Dave Crocker, DEC  
Christian Huitema, INRIA  
Barry Leiner,  
Jon Postel, ISI

## Group 3: SECURITY ARCHITECTURE

Vint Cerf, CNRI [Chair]  
Steve Crocker, TIS  
Steve Kent, BBN  
Paul Mockapetris, DARPA

## Group 4: TRAFFIC CONTROL AND STATE

Robert Braden, ISI [Chair]  
Chuck Davin, MIT  
Dave Mills, University of Delaware  
Claudio Topolcic, CNRI

## Group 5: ADVANCED APPLICATIONS

Russ Hobby, UCDavis [Chair]  
Dave Borman, Cray Research  
Cliff Lynch, University of California  
Joyce K. Reynolds, ISI  
Bruce Schatz, University of Arizona  
Mike Schwartz, University of Colorado  
Greg Vaudreuil, CNRI.

## Security Considerations

Security issues are discussed in Section 4.

## Authors' Addresses

David D. Clark  
Massachusetts Institute of Technology  
Laboratory for Computer Science  
545 Main Street  
Cambridge, MA 02139

Phone: (617) 253-6003  
EMail: ddc@LCS.MIT.EDU

Vinton G. Cerf  
Corporation for National Research Initiatives  
1895 Preston White Drive, Suite 100  
Reston, VA 22091

Phone: (703) 620-8990  
EMail: vcerf@nri.reston.va.us

Lyman A. Chapin  
Bolt, Beranek & Newman  
Mail Stop 20/5b  
150 Cambridge Park Drive  
Cambridge, MA 02140

Phone: (617) 873-3133  
EMail: lyman@BBN.COM

Robert Braden  
USC/Information Sciences Institute  
4676 Admiralty Way  
Marina del Rey, CA 90292

Phone: (310) 822-1511  
EMail: braden@isi.edu

Russell Hobby  
University of California  
Computing Services  
Davis, CA 95616

Phone: (916) 752-0236  
EMail: rdhobby@ucdavis.edu