

Aerospace, Programming Languages, and Information Technology Co-Development: ARPANET and Networking Origins

Brendan Sechter

2026-07-19

This seventh article of the twelve-part series covers the origins of computer networking with focus on the Advanced Research Projects Agency Network hereafter ARPANET as the paradigmatic case. The ARPANET's design and construction between 1966 and 1972 established the engineering foundations on which essentially all subsequent computer networks were built, including the contemporary internet. The role of aerospace and defense research in initiating and funding the network, the technical decisions made during its design, and the transition from a research network to a general-purpose communications infrastructure together illustrate the aerospace-computing coupling formalized in [A237](#) operating on the networking axis of the six-axis framework.

Three technical foundations enabled the ARPANET. Packet switching, developed independently by Paul Baran at the RAND Corporation and Donald Davies at the National Physical Laboratory in the United Kingdom during the early 1960s, replaced the circuit-switching paradigm of telephone networks with a store-and-forward architecture that used network capacity efficiently and survived link failures gracefully. Queueing-theoretic analysis by Leonard Kleinrock at MIT, treated in his 1961 doctoral thesis previously cited in [A237](#), established the mathematical framework for reasoning about network performance under load. Interactive computing terminals and timesharing systems, whose Whirlwind and SAGE origins are treated in [A240](#), supplied the user-facing endpoints that made a research computer network operationally useful. The ARPANET combined these three foundations into a working demonstration that later networks extended without fundamentally displacing.

The Advanced Research Projects Agency

The Advanced Research Projects Agency was established by the United States Department of Defense in February 1958 as a direct response to the Soviet launch of Sputnik 1 on 4 October 1957. The institutional purpose was to prevent the United States from being surprised again by foreign technical achievement, and the organizational mechanism was to fund high-risk high-payoff research at universities and industrial laboratories under the direction of program managers with substantial technical autonomy. ARPA was renamed to Defense Advanced Research Projects Agency hereafter DARPA in 1972 and back to ARPA in 1993, and to DARPA again in 1996. The current name is DARPA, but the ARPANET-era name was ARPA, which is used throughout this article for consistency with the historical record.

The role of ARPA in computing research through the 1960s and 1970s was disproportionate to its budget. The Information Processing Techniques Office hereafter IPTO within ARPA, established in 1962 under Joseph Carl Robnett Licklider, funded essentially all of the important computing research programs of the period including the Compatible Time-Sharing System at MIT treated in [A240](#), the Multics project at MIT and General Electric described in the

primary introduction paper by [Corbató and Vyssotsky 1965](#), the Berkeley timesharing system that later became the foundation of the Unix ecosystem, the Xerox Palo Alto Research Center graphics and personal computing work of the 1970s treated in [Hiltzik 1999](#), and the ARPANET itself. The IPTO culture of long-term high-risk research funding without near-term deliverables produced the computing infrastructure that the commercial computing industry later built on, per the historical treatment in [Norberg and O’Neill 1996](#).

Licklider’s 1960 paper on human-computer symbiosis in [Licklider 1960](#) and the 1968 paper on the computer as a communication device by [Licklider and Taylor 1968](#) laid out the vision of interactive networked computing that ARPANET was designed to realize. The vision emphasized the value of human-computer interaction for augmenting human capability rather than for replacing human labor, and the value of computer-mediated communication for connecting geographically distributed research communities. Both papers are load-bearing primary sources on the intellectual origins of the ARPANET and, by extension, of the contemporary internet.

Packet Switching

Packet switching is the technical foundation that distinguishes computer networks from telephone networks. In a circuit-switched network, an end-to-end path through the network is reserved for the duration of a communication session, with the reserved capacity unavailable to other users regardless of whether it is being actively transmitted on. In a packet-switched network, communications are broken into discrete packets that are individually routed through the network, sharing intermediate link capacity with packets belonging to other communications. The efficiency advantage of packet switching for bursty traffic, meaning traffic patterns dominated by short bursts of intense transmission separated by long idle intervals, is substantial for the interactive and file-transfer traffic that dominates computer network use.

The concept was developed independently by two research groups. [Baran 1964](#) at the RAND Corporation, previously cited in [A237](#), developed packet switching as a survivable communications architecture for the defense problem of maintaining command and control under nuclear attack. Baran’s design used a distributed network of routing nodes each connected to several others, with messages broken into small blocks that could take independent paths and be reassembled at the destination. The survivability property was that no single node or link failure could disconnect the network, and that the network would gracefully degrade rather than fail catastrophically under substantial damage. [Davies 1966](#) at the National Physical Laboratory in Teddington, United Kingdom, independently developed similar ideas for the problem of efficient computer-network communication and coined the term “packet” that became standard in the field.

The efficiency argument for packet switching over circuit switching for bursty traffic can be quantified by the utilization ratio. For a source with peak transmission rate r_p and mean transmission rate r_m across an observation window with utilization fraction $u = r_m/r_p$, a dedicated circuit provisioned at the peak rate carries useful traffic only during the fraction u of the time. Multiplexing N independent bursty sources onto a shared link of capacity

$$B_{\text{shared}} \approx N \cdot r_m + k\sqrt{N \cdot r_m \cdot r_p}$$

for a small safety factor k of order 3 to 5 based on the Central Limit Theorem approximation to independent-source aggregation, provides adequate service to all N sources at a total capacity substantially below $N \cdot r_p$. For interactive traffic with utilization ratios of a few percent, the effective capacity gain from packet switching relative to circuit switching approaches an order of magnitude or more. This efficiency argument was the technical rationale for choosing packet switching in the ARPANET design.

ARPANET Design and Deployment

The ARPANET design was led by Lawrence Roberts, who joined the IPTO in 1966 and became its director in 1969. Roberts drew on Kleinrock’s queueing theory, Baran’s survivability analysis, and Davies’s efficiency analysis to specify a network architecture that would connect approximately 15 to 20 research computing sites at United States universities and defense laboratories. The initial ARPANET plan was presented at the 1967 ACM Symposium on Operating Systems Principles in [Roberts 1967](#), with the subsequent detailed account of the operational network appearing in the [Roberts and Wessler 1970](#) AFIPS paper previously cited in [A237](#). The design decision that Roberts adopted at Wesley Clark’s suggestion was to use dedicated switching computers called Interface Message Processors hereafter IMPs to handle the network communications, with the host computers at each site relieved of the engineering burden of implementing the network protocols. This separation between hosts and switches remains foundational to essentially all subsequent network architectures.

The IMPs were built by Bolt Beranek and Newman hereafter BBN in Cambridge, Massachusetts, under contract to ARPA following a 1968 request for proposals. The primary technical description of the IMP is [Heart Kahn Ornstein Crowther Walden 1970](#), authored by the BBN engineering team responsible for the machine. The IMP was based on the Honeywell DDP-516 minicomputer with special-purpose interfaces to the local host computer and to the leased telephone lines that formed the backbone links. The initial IMP configuration provided approximately 16 kilobits per second per backbone link, with store-and-forward latency across the IMP of approximately 10 milliseconds. End-to-end latency across an N_{hops} -hop path summed the per-hop transmission time and processing time as

$$T_{\text{e2e}} = N_{\text{hops}} \cdot \left(\frac{L_{\text{packet}}}{B_{\text{link}}} + T_{\text{proc}} \right) + T_{\text{propagation}}$$

for packet size L_{packet} , link bandwidth B_{link} , per-IMP processing time T_{proc} , and total propagation delay $T_{\text{propagation}}$ across the leased-line physical distance. Coast-to-coast paths across the ARPANET typically involved four to six IMP hops and accumulated end-to-end latency of approximately 100 milliseconds.

The first four ARPANET nodes were installed between September and December 1969 at the University of California at Los Angeles, the Stanford Research Institute hereafter SRI, the University of California at Santa Barbara, and the University of Utah. The first host-to-host communication attempt was made on 29 October 1969 between UCLA and SRI, with a partial message “LO” successfully transmitted before the SRI host crashed on receipt of the third character, per the retrospective account by Kleinrock as the UCLA principal investigator in [Kleinrock 2010](#) in IEEE Communications Magazine. The intended full message was “LOGIN”, which would have initiated a remote login session from UCLA to the SRI host. The successful full transmission occurred approximately an hour later after the SRI system recovered. This event is generally treated as the operational birth of the ARPANET, though the technical achievement was modest compared with what followed.

Growth and Protocols

The ARPANET grew from four nodes in December 1969 to approximately 15 nodes by early 1971, approximately 40 nodes by early 1973, and approximately 60 nodes by 1975, following approximately exponential growth of the form

$$N_{\text{nodes}}(t) \approx N_0 \cdot e^{t/\tau}$$

with time constant τ of approximately 15 to 18 months over the 1969 through 1980 period, similar to the transistor-density doubling time treated in the substrate section of [A237](#) and consistent with the general pattern of exponential adoption growth for successful network technologies. The first host-to-host protocol was the Network Control Program hereafter NCP, standardized in 1970 and used through 1982. NCP handled the problem of establishing sessions between programs running on different hosts and reliably delivering byte streams between them. NCP was designed under the assumption that the underlying IMP network would deliver packets reliably and in order, which was true for the initial ARPANET architecture but did not generalize to interconnected networks of the kind that emerged in the mid-1970s.

The Transmission Control Protocol hereafter TCP was designed by Vinton Cerf and Robert Kahn between 1973 and 1974 to solve the problem of interconnecting the ARPANET with packet radio networks, satellite networks, and other emerging networks whose delivery guarantees differed from the ARPANET's. The primary description in [Cerf and Kahn 1974](#) in the IEEE Transactions on Communications introduced the split between a lower-layer datagram protocol that provided best-effort delivery and an upper-layer stream protocol that provided reliable in-order delivery, laying the foundation for what later became the split between Internet Protocol hereafter IP and TCP. The TCP/IP split was formalized in 1978 and standardized in [Postel 1981](#) as Request for Comments 793, which became the standard by which the ARPANET, packet radio networks, satellite networks, and subsequent networks interconnected.

TCP's sliding-window flow-control mechanism bounded the effective throughput on any single connection by the window size divided by the round-trip time,

$$B_{\text{effective}} = \frac{W}{T_{\text{RTT}}}$$

which for the 16-kilobit-per-second ARPANET backbone links and 100-millisecond coast-to-coast round-trip times required window sizes on the order of 200 bytes to fill the link capacity. Retransmission of lost or corrupted packets required an estimate of the round-trip time to determine when to declare a packet lost, later formalized by [Jacobson 1988](#) as the adaptive estimator

$$T_{\text{RTO}} = \bar{T}_{\text{RTT}} + 4\sigma_{\text{RTT}}$$

using exponentially weighted moving averages of the mean and standard deviation of observed round-trip times. The engineering discipline of retransmission-timer estimation became a load-bearing concern of network protocol design and remains a subject of active research in contemporary networks.

The value of a network to its users grows more rapidly than the number of users because each new user gains access to all previous users and each previous user gains a new potential communication partner. This scaling, later formalized as Metcalfe's Law in the context of Ethernet local networks per [Metcalfe and Boggs 1976](#), gives a network value of order

$$V_{\text{network}} \propto N(N-1)/2 \approx N^2/2$$

for N connected users under the assumption that all pairwise connections carry equal value. The number is disputable in practice because not all pairwise connections carry equal value, but the quadratic scaling captures the important qualitative property that network value grows faster than linear in the user count and produces the incentive for networks to become universal rather than to segment into disconnected islands.

Aerospace Research Community and the ARPANET

The ARPANET was funded by ARPA, a defense research agency, and its initial nodes were at universities and industrial laboratories doing defense-related research. The set of institutions connected in the first several years included MIT, Stanford, Utah, UCLA, and Carnegie Mellon among the universities, and RAND, MITRE, and BBN among the industrial and non-profit laboratories. Aerospace research communities were prominent in the institutional participation. Documented early MIT ARPANET nodes included MIT Multics, the MIT Artificial Intelligence Laboratory, and MIT Project MAC, though the MIT Instrumentation Laboratory treated in [A242](#) was primarily served by these adjacent MIT nodes rather than by a directly attributed Instrumentation Laboratory host. The Stanford Artificial Intelligence Laboratory, Carnegie Mellon's computer science department, and the various NASA-funded university research groups all participated actively.

The value of the ARPANET to the aerospace research community was that it enabled resource sharing across geographic distances that would otherwise have required physical travel or file transfer by mail. A researcher at Stanford could log into a specialized simulation computer at Utah, transfer results to a display terminal on an MIT timesharing system, and exchange messages with a collaborator at NASA Ames Research Center in Mountain View, all within a single research session. The efficiency gain relative to the pre-network research process, per the account in [Abbate 1999](#) and in [Hafner and Lyon 1996](#), was substantial and drove sustained institutional demand for network capacity that outstripped ARPA's initial expectations.

The ARPANET was transferred from ARPA to the Defense Communications Agency in 1975 as it transitioned from a research network to an operational one, and split into the ARPANET research network and the MILNET operational military network in 1983. The MILNET carried operational military traffic including logistics, personnel management, and command-and-control functions. The ARPANET research network continued to grow and interconnect with other networks including the National Science Foundation Network hereafter NSFNET. The final ARPANET decommissioning in 1990 marked the transition of the research-network traffic to the NSFNET, which in turn transitioned to the commercial internet in the mid-1990s.

Framework Application to the ARPANET

The six-axis framework introduced in [A237](#) applies to the ARPANET with axis weightings reflecting the character of networking as a computing thread.

The first axis is numerical computation demand. Network protocols themselves imposed modest computation demand relative to the applications running over them. IMP processing consumed a small fraction of the Honeywell DDP-516 minicomputer's capacity. Host protocol implementations consumed a small fraction of the host computer's capacity. The computational innovation of network protocols was the algorithmic design of retransmission, flow control, congestion avoidance, and routing rather than the arithmetic throughput.

The second axis is real-time control. Networking was not real-time in the sense of aerospace flight control, but network protocols imposed timing constraints on their own operation. Retransmission timers, round-trip time estimators, and congestion-window updates all depended on timing measurements that had to be maintained within bounds to preserve protocol correctness. The timing engineering practice for network protocols drew on real-time operating system techniques treated in [A241](#) and contributed back to those techniques through the operational experience with distributed timing coordination.

The third axis is reliability and verification. The ARPANET achieved high reliability through the architectural decisions of packet switching, redundant path routing, and end-to-end retransmission of lost or corrupted packets. For a single physical path of N_{links} links each with per-link availability R_{link} and independent-failure assumption, the path availability is

$$R_{\text{path}} = R_{\text{link}}^{N_{\text{links}}}$$

which for high per-link availability approaches unity for small N but declines rapidly for large N . Redundant path routing across k independent paths raises the aggregate availability to $1 - (1 - R_{\text{path}})^k$, which is the quantitative argument for the packet-switching survivability property that motivated Baran’s original design. Verification of network protocols became a distinct engineering discipline in the 1980s and 1990s as network scale grew and as security concerns emerged.

The fourth axis is networking and distribution. The ARPANET is the paradigmatic case for this axis. Essentially all of the engineering practices that define computer networking as a discipline originated in or were substantially shaped by ARPANET experience. IMP queue occupancy at steady state followed Little’s Law relating the mean number of packets in the queue to the arrival rate and mean queueing time,

$$L_{\text{queue}} = \lambda \cdot T_{\text{wait}}$$

for packet arrival rate λ and mean wait time T_{wait} , which combined with Kleinrock’s queueing-theoretic analysis cited in [A237](#) gave the operational analysis framework for IMP sizing and network capacity planning. The transition from research network to operational network to commercial internet illustrates the general pattern of aerospace and defense computing giving rise to commercial computing infrastructure treated throughout the series.

The fifth axis is software engineering as a discipline. Network protocol implementation drew on and contributed to the software engineering practices treated in [A240](#) for large-scale systems. The requirement that protocols work correctly across independent implementations by multiple developers at multiple sites produced the engineering discipline of interoperability testing, protocol conformance evaluation, and specification refinement that later spread throughout software engineering.

The sixth axis is semiconductor economics and dual-use. The ARPANET drove demand for the class of minicomputers that BBN used for the IMPs and that host sites used for network implementations. The commercial minicomputer industry that Digital Equipment Corporation and its competitors built treated in [A240](#) served both the ARPANET and the broader commercial market, with the network-computing applications contributing to the sustained demand that kept the minicomputer industry viable through the 1970s.

Conclusion

The ARPANET established the engineering foundations on which essentially all subsequent computer networks were built. Packet switching supplied the architectural principle. Kleinrock’s queueing theory supplied the mathematical framework. The IMPs supplied the mechanism for separating network communication from host computation. TCP/IP supplied the interconnection protocol that let the ARPANET connect to other networks and eventually to the commercial internet. The institutional support from ARPA supplied the funding continuity that permitted long-term research investment without near-term commercial pressure. The aerospace and defense research community that supplied the initial user base drove the requirements that shaped the network design, and the pattern of aerospace-funded research producing commercial infrastructure illustrates the general dual-use spillover mechanism treated in [A237](#).

The next article in the series treats the Space Shuttle primary avionics software system as the first large-scale demonstration that safety-critical software could be built to airline-

comparable reliability targets, HAL/S as a purpose-built aerospace programming language, and the IBM Federal Systems Division software process as a template for later programs.

References

Books

- [Abbate 1999](#)
- [Hafner and Lyon 1996](#)
- [Hiltzik 1999](#)
- [Norberg and O'Neill 1996](#)

Reference

- [ARPA/DARPA History](#)
- [Internet Society History of the Internet](#)

Related Posts

- [A237 Framing and the Co-Development Mechanism](#)
- [A240 Early Cold War Air Defense and SAGE](#)
- [A241 Aerospace Simulation and Real-Time Systems](#)
- [A242 The Apollo Guidance Computer](#)

Research

- [Baran 1964](#)
- [Cerf and Kahn 1974](#)
- [Corbató and Vyssotsky 1965](#)
- [Davies 1966](#)
- [Heart Kahn Ornstein Crowther Walden 1970](#)
- [Jacobson 1988](#)
- [Kleinrock 2010](#)
- [Licklider 1960](#)
- [Licklider and Taylor 1968](#)
- [Metcalf and Boggs 1976](#)
- [Postel 1981](#)
- [Roberts 1967](#)
- [Roberts and Wessler 1970](#)