

Ultracomputers: A Teraflop Before Its Time

Gordon Bell

Computer designers have been striving for a decade to build supercomputers that run at speeds near one teraflop (10^{12} floating point operations per second). Accelerating this achievement would require the development of what I term "ultracomputers" (1) that heavily rely on parallel processing. Already, first-generation ultracomputers are available, consisting of thousands of networked computers and costing between \$50 million to \$300 million (Fig. 1). But these machines yield high performance only in specialized, highly parallel applications, and this in turn requires new algorithms and software. In my judgment, substantially more powerful computers will be available in 1995 that will offer teraflop performance for the cost of present supercomputers. Work in progress and developments on the horizon promise an era of "commodity supercomputing." Better computers will be available in 1995 if the government funding that would be wasted on purchasing present ultracomputers were turned instead toward training and software to exploit the power of the next generation.

In 1989, I described the situation in high-performance computing in science and engineering, and specifically mentioned several parallel architectures that could deliver teraflop power by 1995, assuming no constraints on price (2). My prediction was that either of two alternatives could achieve this

goal: (i) thousands of processing elements, each operating on its own data stream, controlled by a single instruction stream (SIMD) or (ii) multicomputers with over a thousand interconnected, independent computers. The sharing of memory among several processors did not look feasible, and I suggested that traditional supercomputers like the Cray, with multiple vector processor architecture, would not evolve to a teraflop until the year 2000.

Compare this with what has occurred since my predictions: During 1992, NEC's 4-processor SX3 became the fastest computer (3), delivering 90% of its peak 25.6 gigaflops for the LINPACK benchmark (a set of numerical calculations), and Cray's 16-processor C90 provided the greatest throughput for supercomputing work loads. Traditional supercomputers deliver approximately 600 to 1,000 flops per dollar. Also at this time, the SIMD approach was abandoned by Thinking Machines, Inc., because it was only suitable for a few, very large scale problems and uneconomical for typical computing work loads.

Moreover, Intel and Thinking Machines introduced massively parallel multicomputers (mmC) based on 32-bit "Killer" CMOS (complementary metal-oxide semiconductor) processors. In 1992, CMOS microprocessors deliver 5,000 flops per dollar—with the current rate of processing, 25,000 flops per dollar will be available in 1995. These new designs join multicomputers from Alliant, AT&T, IBM, Meiko, Mercury, NCUBE, Parsytec, Transtech, among others, and Convex, Cray, Fujitsu, IBM, and NEC are all working on new generation 64-bit massively parallel multicomputers. By 1995, it appears this large number of efforts, together with the evolution of fast, local-area network-connected workstations and "killer" CMOS processors, will usher in an era of commodity supercomputing.

Another development has been the introduction of the KSR-1 shared memory massively parallel multiprocessor by Kendall Square Research. This architecture uses 1,088 64-bit microprocessors tied together

through a distributed memory scheme called ALLCACHE, which eliminates physical memory addressing. Work is not bound to a particular memory location, but moves to the processors that require the data. This approach is flexible, because any processor can be deployed on either scalar or parallel applications; it is general purpose, equally useful for scientific and commercial processing; and the KSR-1 runs traditional supercomputer FORTRAN programs with high throughput. Running the risks shared by all prognosticators, I will again peer into the future and predict that the KSR architecture—namely, a shared virtual memory multiprocessor—is the most likely blueprint for future massively parallel computers.

Unfortunately, one factor that could distort the natural evolution of supercomputing is government involvement. The teraflop quest is fueled by the massive High Performance Computing and Communications program, and by DARPA's focus on teraflops. Gigabuck programs such as this are certain to accelerate the quest at the expense of programmability, usefulness to a large number of users, and long-term development. Already, the existence of government-sponsored architectures has led to the elimination of benchmarking, open bidding, and widest utility for a narrow focus on the teraflop. Although DARPA has a long and successful record of sponsoring university research that has created products, companies, and even industries, its role in the development of high-performance computers through selecting designs should be ended because it has been picked up by industry.

Whether traditional supercomputers or massively parallel computers provide more computing, measured in gigaflops per month, in 1995 is the subject of a bet between Danny Hillis of Thinking Machines and myself (4). Traditional or "true" supercomputers are likely to supply much of the power this decade because of the installed software base and programming methods. In my view, with a free computing market, devoid of government mandated architectures and where users are free to select the machines they use, the main direction will be the shared memory multiprocessor (5).

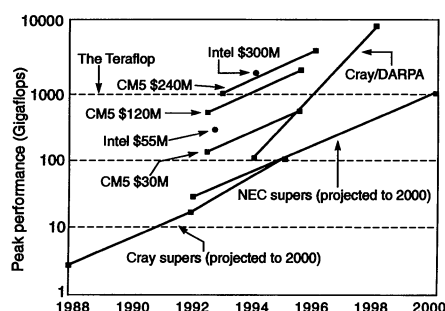


Fig. 1. The race for the teraflop. Peak performance of selected supercomputers and ultracomputers with performance projections. CM5 models from Thinking Machines, Inc., at \$30, \$120, and \$240 million; Intel Paragon models priced at \$50 and \$300 million for 0.3 and 1.8 teraflops; Cray Research supercomputers (Cray YMP/8, C-90, and extrapolated models) priced at \$30 million; NEC supers are the SX3 series supercomputers extrapolated at \$30 million; Cray/DARPA is the performance target for Cray's massively parallel computers for which DARPA has contracted.

The author was manager of research and development for Digital Equipment Corporation from 1961 to 1983 and is a former professor of electrical engineering at Carnegie Mellon University. He has been a technical adviser to Kendall Square Research, Cambridge, MA, since its beginning and a consultant to the Intel Corporation and Thinking Machines.

REFERENCES AND NOTES

1. A more extensive discussion of these issues will appear in an article by G. Bell, *Commun. ACM* (in press, August 1992).
2. G. Bell, *ibid.* 32, 1091 (September 1989).
3. T. Watanabe, *The NEC SX3 Supercomputer* (University Video Communications, Stanford, CA, 1992).
4. J. L. Hennessey and D. A. Patterson, *Computer Architecture: A Quantitative Approach* (Kaufman, San Mateo, CA, 1990).
5. G. Bell, in *Carnegie Mellon Computer Science, 25th Anniversary Commemorative*, R. Rashid, Ed. (Addison-Wesley, Reading, MA, 1991), pp. 3-27.