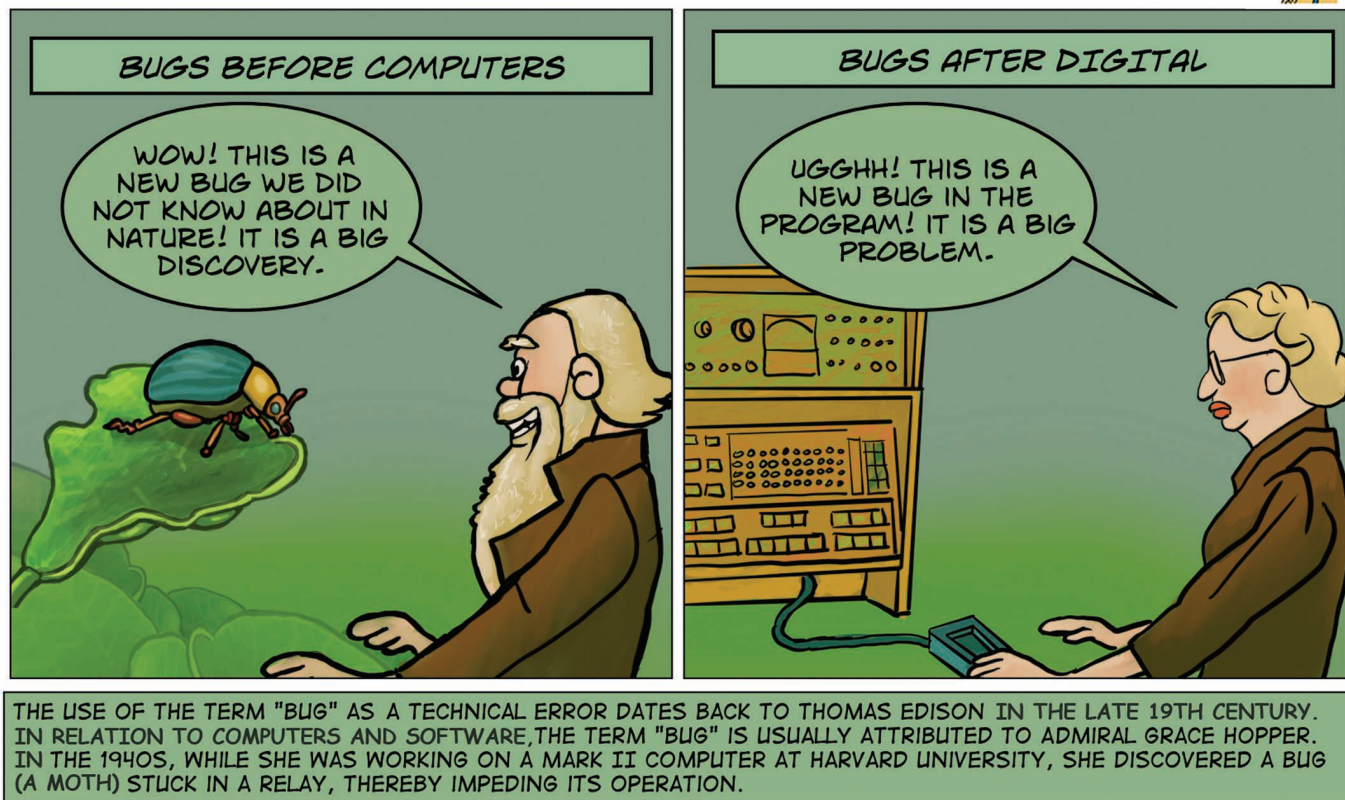


COMPUTING THROUGH TIME



Digital Object Identifier 10.1109/MC.2021.3113774
Date of current version: 17 November 2021

Improving Application Performance on the HP/Convex Exemplar; Thomas Sterling et al. (p. 50) “The Exemplar, built by HP/Convex, is among the newest entries in this class. The Exemplar combines the parallel scalability of MPPs with hardware support for distributed shared memory, global synchronization, and cache-based latency management.” (p. 52) “The applications from the Earth and space sciences project that were used in this work are the piecewise parabolic method (PPM), a finite-element method (FEM) for unstructured meshes, a tree code for the n -body problem (Tree), and a particle-in-cell (PIC) code.” (p. 54) “Ultimately, our experience shows that the uniform name space is a significant enhancement for programmability but that cache-based latency management alone is insufficient to achieve acceptable performance for many problems. Even in the distributed-shared-memory context, programmers must be involved in partitioning data and scheduling tasks to minimize overall latency.” [Editor’s note: The analysis done here for improving the use of SMMP architectures shows what is true even today: the parallelization of algorithms is not an easy task, and it relies on the optimal positioning of many parameters.]

Application Performance on the MIT Alewife Machine; Frederic T. Chong et al. (p. 57) “In this article, we present the performance of 14 applications on the Alewife machine, including both coarse- and fine-grain applications. Not surprisingly, Alewife’s mechanisms support the good performance of traditional coarse-grain applications from the Splash and NAS benchmark suites. But we also show that Alewife provides excellent communication mechanisms for fine-grain applications, even without data reuse.” (p. 64) “The results confirm that hardware support for limited sharing is adequate for a broad range of applications, even on large numbers of processors. Local cache-miss behavior turns out to be important on multiprocessors with low remote miss latencies.” [Editor’s note: This analysis of a wide range of scientific algorithms shows, like other research results, that remote and local cache misses, that is, data distribution schemes, play a more important role than load balancing alone.]

Shared Memory Consistency Models: A Tutorial; Sarita V. Adve et al. (p. 66) “The memory consistency model of a system affects performance, programmability, and portability.