

Mathematical sciences in the nineties

W. R. Pulleyblank

In the last decade of the twentieth century, we saw great progress in the mathematical sciences as well as changes in the activities of the mathematical scientist. These included the resolution of some well-known conjectures, the introduction of new areas of study, as well as the adoption of new tools and new methods of operation. The IBM Research Division was an active participant in many of these events. We discuss here a selection of these, focusing on some to which contributions were made by mathematicians at IBM Research.

Introduction

During the last decade of the twentieth century, the world underwent a number of significant changes. The Cold War, which had dominated so much of our political activities, was over. The Internet emerged from universities and national laboratories to become a dominant part of the business, home, and recreational infrastructure. We saw changes in the investment focus of governments and reprioritization of the use of funds by granting agencies. During this decade, the personal computer evolved from being a “smart” terminal to become a powerful platform for computing, communication, and network management. We saw the emergence of pervasive computing devices. The success of the Palm Pilot showed that people could even be taught a new method of printing, using “graffiti”—a symbolic method that was not unnatural, but which did have some definite rules on how letters were to be formed.

The mathematical sciences experienced exciting developments and changes. Fermat’s Last Theorem was finally, definitively, proved. Even more remarkable, some mathematicians became folk heroes to the general public. The principal characters of the successful play *Copenhagen* may have been more theoretical physicists than mathematicians, but there was no question that *Proof*, which went from being an off-Broadway Tony Award winner to a Broadway hit, was about mathematicians. *A Beautiful Mind* [1], Sylvia Nasar’s biography of John Nash, not only received acclaim itself, but became the basis for the film of the same name, which was awarded the 2002 Academy Award for best picture.

Throughout history, areas of research in the mathematical sciences have experienced periods of expansion and of consolidation. The former occur when a new field or application of mathematics is discovered and the challenge is to see how the existing models can be applied and what new ones must be developed. The latter occur as unified theories are developed which both simplify the mathematical models and clarify their underlying structures. We have seen examples of both of these in this decade.

In 1960, the Mathematical Sciences Department was created as a department of IBM Research, led by Herman Goldstine, the first Director, from 1960 through 1965. Since then, there has been a long record of successes, and many distinguished mathematicians have been members of the department. It has combined significant scientific research with the application of these results to a range of IBM problems, and continues to be recognized as one of the world’s foremost industrial research departments.

From 1995 to 2000, I served as the Director of Mathematical Sciences in the IBM Research Division. During this time, financial mathematics emerged as an active area, particularly the fields of risk measurement and management. There were significant advances in mathematical modeling and in tools used to solve mathematical models, particularly optimization. Remarkable progress in both the fundamental theory and applications of dynamical systems also took place during this period.

In this paper I review some trends and events of this period, focusing on a selection of the activities

which took place in the IBM Research Division. I discuss the increasing role of computers in the mathematical sciences, as well as the advances in speech and natural language processing. I also discuss the developments in data mining and related statistical methods. Finally, I discuss some developments in publication and dissemination of results, due in large part to the evolution of the Internet. I have made no attempt to be comprehensive, and what follows should be viewed as a somewhat personal selection from a much broader collection of work.

Developments in computing

In March 1993, Intel introduced the Pentium** processor, with a clock speed of 66 MHz. This was rated as having six times the performance of the earlier 33-MHz 486 processor. By the end of the decade, desktop computers had reached a speed of 1.5 GHz and the IBM ThinkPad* broke the 1-GHz boundary in 2001. In this time span, disk storage on ThinkPads had increased from less than one gigabyte to more than 15 gigabytes.

The IBM SP/1* was introduced in 1992. By the end of the decade, the SP* had become the world's dominant parallel computing platform. The Advanced Scientific Computing Initiative (ASCI) program of the Department of Energy had been launched and was driving rapid development in high-end computing capability. IBM had successfully installed the ASCI Blue Pacific (3.8 teraflop/s) and ASCI White (12.3 teraflop/s) supercomputers at the Lawrence Livermore National Laboratory in 1998 and 2000 respectively. As of March 2002, ASCI White was still the most powerful supercomputer on the "Top 500" list (www.top500.org), and ASCI Blue Pacific was rated number five.

The effects of these advances on the mathematical sciences were profound and far-reaching. First, it became possible to prepare, manage, and read mathematical documents on portable computers such as a ThinkPad. Second, symbolic mathematics systems, such as Mathematica**, AXIOM**, and MAPLE**, became much more broadly used. These no longer required high-end workstations or servers, but could also run on a laptop computer. Third, powerful computing platforms became broadly available, with the result that applications would now migrate from their traditional homes—universities, national laboratories and industrial research laboratories—to industrial and commercial sites. (At present, more than half the sites on the Top 500 list of supercomputers are commercial.) In May of 1999, IBM launched the Deep Computing Institute with the goal of supporting this transformation. Mathematical scientists had an impressive and rapidly growing set of tools. Moreover, platforms capable of running sophisticated mathematical software became available to a much broader audience.

Another development that dramatically changed the environment of mathematical scientists was the emergence of the Internet. The number of Internet nodes increased from approximately 100,000 in 1989 to more than 50 million in 2000. The World Wide Web was created in 1991, and the MOSAIC** browser appeared in 1993. Over the decade of the nineties, it became common practice for new research papers to be posted on Web sites, available to anyone. Now it was no longer essential to wait for an author to send a reprint or a electronic version of a paper. Anyone could find it and download it. In the past, access to real-world data was always an obstacle for algorithmic researchers. Now, however, many Web sites contain data for problem instances. For example, TSPLIB (www.iwr.uni-heidelberg.de/iwr/comopt/software/TSPLIB95/) has made it possible for a much larger community of researchers not only to stay up to date on progress in solving real-world traveling-salesman problems, but also to download the data and work on specific instances. The Web also made it easy to establish electronic contact with individuals known only by name by using e-mail. At the end of the decade, the Mathematical Sciences Department launched the Common Optimization Interface (COIN) Web site (<http://oss.software.ibm.com/developerworks/opensource/coin/>)—an open-source site which distributes mathematical programming software and is intended to spur the development of open-source software for the operations research community [2].

An additional area of research was becoming increasingly important: the development of high-performance software for mathematical functions. The new computing platforms being developed were capable of outstanding performance; however, software specialized to the architecture of a specific machine was required in order to realize optimal performance. The IBM Engineering and Scientific Subroutine Library had to be updated for each new family of processors. The Mathematical Acceleration Subsystem (MASS) Web site (<http://techsupport.services.ibm.com/server/mass>) was created, distributing free high-performance computing software to run on IBM platforms. In addition, Gustavson [3, 4] showed that the recursion paradigm, a fundamental technique in theoretical computer science, would almost automatically make use of the cache levels of a hierarchical memory system to deliver near-peak performance on a broad range of scientific calculations.

Data mining and very large datasets

One consequence of the developments in computing platforms has been the creation and availability of very large databases. Two general types of analysis are performed on large datasets. The first, more traditional type is *hypothesis verification*. In this case a conjecture is made by a user, and data is analyzed to see whether it

supports the hypothesis. The second type is *discovery-based data mining*, or simply *data mining*. The objective is to find “unexpected” features which exist in these datasets by virtue of the existence of a large set of records. In this case we do not start with a hypothesis. Rather, we require the data mining system to both generate and verify a hypothesis having a prescribed form. For example, a fact such as “males of age at least 70 who drive sports utility vehicles have lower average claims on their auto insurance than other segments of the population” need not be hypothesized, but can be detected by a data mining system by analyzing data covering a sufficiently large sample of the population. See Apte et al. [5] for a broad discussion of activities in this area.

In general, data mining consists of “optimally” fitting a particular model to a set of data. *Clustering* involves partitioning the data into a prescribed number of sets in such a way that each set is close to being homogeneous, and distinguishable from other sets in the partition. In general there are as many different ways to cluster data as there are ways to define the distance between two data subsets. Usually the number of sets or the size of sets in the cluster is constrained so as to avoid trivialities. For example, *protein fold recognition* attempts to partition all known proteins into a small number of classes on the basis of their three-dimensional structure.

Classification is, in some sense, the opposite operation. It involves taking data which has already been (somehow) divided into classes and then attempting to find features in the data that support, or explain, this division. For example, if the complete history of all customers of an insurance company were in a single database, we could divide the records into two groups depending on whether or not a claim had been submitted. Classification would then try to identify *other* features, or combinations of features, in the records which would indicate the class to which a given customer belonged. The goal is to find rules having predictive power, which can subsequently be applied to unclassified data to determine the class into which the records would fall.

Detection of *associations* involves finding data elements which tend to occur together in a large database. For example, if we have a database containing point-of-sale information, we can use association analysis to determine groups of items that tend to be purchased together by a store’s customers.

There are strong connections between data mining methods and statistical methods used in analyzing data. Traditional statistical methods are usually used for analyzing numerical data which comes from some measurement; however, data mining is particularly useful on coded data. In this case, the coded data value may have no inherent importance, but simply indicates a class to which the data belongs. For example, whether a person

has submitted an insurance claim in the past five years is indicated by coded data. However, the number of claims submitted is a numerical entity and may be the subject of statistical analysis. Statistical analysis is an important component of advanced data mining methods. For example, in a study on catalog market saturation, statistical analysis established that the responses from catalog mailings followed a log-normal distribution. This enabled data mining to predict responses to mailings more effectively.

IBM scientists were very active in data mining during the 1990s. Agrawal and colleagues in Almaden formulated the problem of finding associations and also developed efficient algorithms for finding associations [6]. Michaud and colleagues in Paris developed methods for data clustering. Hong [7] showed how algorithms for circuit optimization could be adapted to provide a powerful classification tool, called R-MINI. Apte and colleagues ran successful studies which showed how these data mining tools could be applied and adapted to deal with a broad set of real-world situations. One example is modeling customer response to mailings by a catalog sales company with levels of success equal to or better than those achieved by the company’s market analysis experts (Apte et al. [8, 9]). Pednault developed a system called Probe which combines some statistical features with clustering. For example, it can partition the set of records of customers of an insurance company into groups of similar individuals such that the variance of the amount of claims of the members of each class is small. (See Apte et al. [10].) This has turned out to be a powerful and widely useful tool.

Optimization

The 1990s was a period of rapid development in optimization theory and methods. In the prior decade, the ellipsoid method and barrier methods had been developed for linear programming. These had been shown to have important theoretical consequences. In addition, the latter had led to very effective computational tools for a broad range of real-world problems. In the 1990s, significant developments were made in these codes, and in the well-established simplex algorithm itself. Moreover, major advances were made in integer programming methods. These combined theoretical advances with great improvements in software methods and the underlying platforms. It now became routine to efficiently solve problems one to two orders of magnitude larger than had previously been possible. In addition, solutions to problems of a fixed size could be obtained much more quickly. For example, one airline fleet scheduling problem instance obtained from United Airlines has served as a test benchmark at IBM Research since 1975. At that time, it took approximately eight and a half hours to solve on a

large mainframe. By 1990, new methods, computer codes, and hardware platforms had reduced the running time to ten minutes. By the end of the decade, this problem could be solved in less than a minute on a ThinkPad. This represents an overall computation speed increase by a factor of more than 1500.

This increase in speed has had profound consequences. A problem which formerly had required a full working day to solve on an enterprise-sized machine could now be solved many times per hour on a moderately sized workstation. This has created the possibility of exploring a broad range of alternative scenarios in order to develop robust solutions which will remain valid even if conditions change. It has also created the opportunity for *continual optimization*, wherein we actually reoptimize in real time to respond to changes in the data. The latter subject is of particular interest in the airline industry. This industry has been a leader in using optimization as a strategic and tactical planning tool, but has only recently been able to use it as part of its operational procedures.

Commercially available optimization software became more powerful, more reliable, and more broadly used. For example, in the mid-nineties, the IBM Research Division, together with Sabre Decision Technologies, developed a crew-pairing optimization system that was used operationally by four different airline companies. It enabled airlines to develop schedules for assigning crews to aircraft that significantly reduced their operating costs. This system used the Optimization Subroutine Library (OSL), the premier IBM optimization product, as an engine. It also contained extensive software for managing libraries of solutions and for manually modifying solutions to better deal with “intangible” requirements. See Anbil et al. [11].

The nineties also saw the development and deployment of effective heuristic methods for problems for which no general solution methods are known. These typically would perform well and had the additional advantage that they would yield feasible solutions to complex problems, solutions that satisfied all of the constraints of a problem. An example of this was the scheduling of maintenance personnel work performed by Bar-Noy et al. [12, 13], which resulted in personnel-scheduling software currently in use by IBM Global Services.

Theoretical computer science has focused on the fundamental questions of finding more efficient algorithms for problems as well as establishing lower bounds on the work required to solve problems. During the nineties, we saw major advances in so-called *approximation* algorithms. Prior to this, most research had concentrated on determining the complexity of computing optimal solutions to problems. Now the attention shifted to efficiently finding solutions which might not be optimal, but which could be guaranteed to come within a certain percentage

of optimality. For example, it was well known that solving a maximum cut problem in a weighted undirected graph was NP-complete.¹ But David Williamson of the IBM Research Division and Michel Goemans of MIT [14] developed a polynomial bounded method which would compute a solution whose value was no worse than 0.878 of the optimum.

Another development in the theory of computation that took place in the nineties was the establishment of the area now known as “the foundations of computational mathematics,” or FOCM. This began by establishing a theory for the limits of computation over the real numbers, independent of their binary representations. Subsequently it was broadened to include many areas of computation. There is now a journal entitled *Foundations of Computational Mathematics* (edited by Michael Shub of the Mathematical Sciences Department), and major biennial conferences are held in this rapidly developing field. See [15] for an editorial by Shub discussing the goals of this activity.

Speech recognition, natural language understanding, and text mining

Another area of research, with strong mathematical connections, which flourished in the nineties involved the processing and analysis of natural language. Work on speech recognition, which had been proceeding for more than twenty years in the IBM Research Division, made significant advances. It was now possible to recognize and interpret human speech without requiring spaces between the words, limited domains of discourse, or extensive customization for individual speakers. The “name dialer” is in operational use at IBM: When a caller speaks the name of the desired party, the system will understand it and find the phone number. In addition, the ViaVoice* products have evolved as market leaders in the area of speech dictation.

The methods used depend heavily on statistical methods. The goal of automatic speech recognition is to transform a given speech utterance into a text of word sequence that best matches the speaker’s intended word sequence. The approach used starts with a probabilistic formulation of speech recognition as a noisy communication channel problem. Techniques such as hidden Markov models, decision trees, and Gaussian mixtures are used as modeling tools. Statistical estimation techniques such as maximum likelihood and maximum mutual information are used to estimate the parameters of the speech recognition probabilistic model. See Bahl et al. [16] and Jelinek [17].

Natural language understanding can be applied to the text obtained from speech recognition or directly to text

¹ See <http://mathworld.wolfram.com/NP-CompleteProblem.html>.

obtained from other sources. The goal of *text mining* is to discover the content that the creator of the text had in mind when it was produced. There are several approaches. One is to use partial grammars which enable parsing of the text and extraction of meaning from the resulting structure information. Another is to apply statistical techniques analogous to those used in speech recognition. Another is to base the analysis on the presence or absence of certain keywords.

Research has been carried out using all of these approaches at IBM, with a steady stream of results. Impressive demonstrations have been created enabling telephone banking, management of mutual fund portfolios, and booking of airline flights. The Web Genie system of David Johnson and colleagues enabled a natural language interface with a Web site which simplified the search for information at a financial Web site. This extended earlier work on automatic classification of e-mail. See Johnson et al. [18].

Interestingly, the best systems for speech recognition make use of natural language understanding to ensure that the result “makes sense.” If it does not, alternative interpretations of words are tried in an attempt to find a result that is consistent with the utterance and seems to have some real meaning.

Dynamical systems and applications

Dynamical systems provide a framework for modeling and understanding a broad variety of phenomena related to the evolution of systems. In particular, understanding chaotic dynamics is an important problem—first, as a key to understanding deep properties about instabilities and chaos in natural sciences, and second, as a tool for practical applications ranging from mechanical engineering to digital printing.

The transition to chaos plays a central role in the theory of dynamical systems. Bifurcation theory describes the mechanics of these transitions. In the mid-seventies, Coulet and Tresser, and, independently, Feigenbaum, adapted renormalization group techniques to the description of simple models of turbulence. They discovered surprising universal constraints on the geometry at the time of transition to chaos.

Renormalization was a tool introduced to understand this universality. During the nineties, research carried out in the department showed how to formulate renormalization in a variety of contexts in one dimension and beyond and unveiled surprising relations linking smoothness, dimension, topology, and universality in more realistic models of nature. See Martens and Nowicki [19].

Shub and collaborators made a breakthrough in this time period by showing that large classes of dynamical systems exhibit robust statistical behavior. In order to prove these results, Pugh and Shub [20] developed new

concepts in nonlinear dynamics, namely Julienne quasi-conformality and holonomy density point preservation. This work will almost certainly have significant long-term impact on the development of the mathematical theory of dynamical systems. See Burns et al. [21].

Applications occurred in several domains. Work which had begun earlier on recording technology and channel encoding continued into the nineties, increasing the competitiveness of IBM storage devices. During this period, one of the more surprising new applications was the application of dynamical systems to digital halftoning. This is the problem of representing a range of shades of gray by an appropriate pattern of dots in such a way that the fidelity of the original picture is maintained while avoiding the introduction of distracting “artifacts.” Work began on black-and-white high-end printers and continued on the IBM top-of-the-line color printers. See Adler et al. [22] for a description of the applications to printing.

Financial mathematics, risk measurement, and risk management

The development of the theory of probability as a companion to statistics is well chronicled in Bernstein [23]. In the nineties there was considerable activity in the fields of risk measurement and management, particularly as it related to the mathematics of finance. This was very interesting for several reasons. First, risk itself became a concept to be quantified and studied. It became understood that not only was it possible in some cases to reduce risk by appropriate choices of policy, but it was also possible to assign a value to the reduction of risk and hence pay an individual to accept risk or, conversely, to accept payment to assume someone else’s risk. At present, options have become common financial instruments. An option basically gives the right to buy or sell goods at a prescribed price at some point in the future. An option has a price based on the value of the goods as well as the perceived volatility of this price. In 1997, the Nobel Prize for Economics was awarded to Robert C. Merton and Myron S. Scholes for their work “for a new method to determine the value of derivatives,” which led to a method of computing these values in certain circumstances.

A second reason why this subject is of great interest is that it provides an example of a phenomenon that can be modeled which is not based on properties of physical objects. For that reason, established mathematical models often do not deal well with these situations, and new types of mathematics must be developed to deal with them.

Three broad classes of risk are studied at present. The first, *market risk*, attempts to determine the uncertainty in price of an object traded in a liquid market. The second, *credit risk*, attempts to place a value on the uncertainty associated with an account receivable. How should we account for the possibility that a debtor may default

on an obligation? The third, *operational risk*, basically tries to handle everything else. It considers the full set of other risks that a business must face, including risk of catastrophic political events, weather-related risk, and risk of criminal activity. Not surprisingly, this is the least well understood class of risks. Operational risk is beginning to receive attention. It is difficult to manage because it focuses on rare events which require new methods of treatment. Market risk is the most extensively studied, and considerable mathematical work has been done here. (See Crouhy et al. [24] for an extensive treatment of this subject.) One aspect of market risk was treated by Xin Guo [25], who showed how the now-classical Black-Scholes model could be extended to value insider information in a market.

The basic elements of study are individual financial instruments—for example, stocks, bonds, and mortgages. Much modern investment theory centers around portfolios of investments. Intuitively, a portfolio reduces risk by diversification, but it also has a lower expected future value than if all investments were put into the single instrument with the highest expected future. In 1990, the Nobel Prize for Economics was awarded to Harry M. Markowitz (who was a member of the IBM Research Division from 1974 to 1983) “for pioneering work in the theory of financial economics.” Markowitz showed how to use the variance of the value of a portfolio to represent risk and how to optimize the expected value of a portfolio, given an acceptable risk threshold, or conversely, how to minimize the risk required to obtain a desired expected return.

Jensen and King [26] showed how to extend this theory to handle a much broader set of nonsymmetric risk functions. In addition, Jensen (unpublished) showed how to use quadratic integer programming, present in OSL described above, to extend the Markowitz theory and produce the first system that could recommend portfolio transactions responding to changes in market dynamics, taking into account transaction costs, minimum desired positions, discreteness of order sizes, and limits on investments within certain sectors.

Publication and the mathematical sciences

The publication of results in refereed archival journals has been a primary channel for output of mathematical results. It is a massive self-managed enterprise, and is probably one of the most effective examples of knowledge management. The editors are recognized scholars with distinguished publication records. The referees are themselves experienced researchers in the relevant topics, capable of assessing the correctness, importance, and originality of submitted results. The result has been the creation of recognized top-tier journals with well-earned reputations for publishing significant results with long-

lasting value. Members of the Mathematical Sciences Department have served as editors for many of these journals.

The publishing companies themselves have typeset the papers, produced and distributed the journals, and promoted their publications. This provided relatively broad distribution of the papers, in part through individual subscriptions, but mainly through library subscriptions. A researcher could easily and quickly obtain copies of papers in these journals just by knowing the bibliographic reference information. One continuing problem has been the frequent long delays between submission of a paper by an author and its subsequent appearance. Reducing these delays has been a major goal of the editorial boards of almost all major journals. To avoid these delays, it became commonplace for authors to distribute preprints of papers to colleagues whom they knew were working on similar topics.

The Internet and broad availability of computing platforms is substantially changing this process. Now preprints are almost always distributed electronically. Moreover, most researchers post these preprints on Web sites so that they are accessible by anyone who wants them and has access to a search engine. Typesetting systems such as TeX and LaTeX produce results which may not reach the level of the typesetter’s art seen in papers produced by professionals, but they are equally effective in communicating the results. The IBM techexplorer Hypermedia Browser* is a sophisticated viewer created by IBM for these types of documents which lets them be viewed from standard Web browsers. Adobe Acrobat** allows facsimiles of documents to be distributed over the Web and protects the document from modification by a user. (See <http://www.ibm.com/software/network/techexplorer/>.)

The access to information for a researcher is immeasurably better than it was before. Now researchers in any country of the world can have equal access to material as soon as it becomes available, provided that they have Internet access. This is independent of whether it has “appeared in print.” Moreover, as soon as a document is recorded digitally, it can be managed using a broad set of tools. Search engines can be much more effective than conventional indices or tables of contents when looking for information in a paper, or for a paper containing certain information. Reproduction and dissemination of the paper itself is much simpler. Corrections can be made to a paper on a Web site, and from then on, all researchers who access the paper will see the new version.

Publishing companies are facing significant challenges. As libraries come under increasingly severe budgetary restrictions, they are reducing the number of journals to which they subscribe. The cost of the journal is often a major factor in the decision to keep or terminate a

subscription. Many journals are finding their institutional subscriptions decreasing, and they have not yet been able to develop an economic model for electronic journals, although many are trying. Increasingly, their distribution and production functions are being augmented or replaced by individual researchers on networked computers.

Refereed journals and conference proceedings do offer several real benefits. One enduring value of traditional publication is the endorsement of the correctness of the material in the paper by the editorial board. Another is its scholarly quality. This includes appropriate references to previously published art, standard terminology, and quality of exposition. Another is the permanence of a printed paper. Electronically recorded information must continually be converted to new formats and re-archived to reflect changes in computing hardware. (Anyone who still has data on 5¼-inch diskettes understands this issue.)

We are already seeing a shift in some disciplines, such as computer science, toward rapid publication of unrefereed, or superficially refereed, conference proceedings. I believe that it will still be possible for publishing companies to succeed in this new electronic world, but it will require rapid adaptation to accommodate the emerging styles of research and electronic communication of researchers. The challenges faced by scientific publishers will, I believe, be no less daunting than those currently faced by record and film companies as they attempt to protect and manage their content while adapting to a new world of e-commerce.

Conclusion

The Department of Mathematical Sciences in the IBM Research Division has been and continues to be a top-quality academic and industrial research unit. The research agenda continues to evolve, maintaining its relevance to the IBM Corporation and its connections with the worldwide mathematical sciences community. As new application areas for mathematics develop, and as existing areas evolve, it can continue to play a broad and central role at IBM Research and an influential role in global mathematical sciences activities as we move into the 21st century.

*Trademark or registered trademark of International Business Machines Corporation.

**Trademark or registered trademark of Intel Corporation, Wolfram Research, Inc., Numerical Algorithms Group, Waterloo Maple Inc., Netscape Communications Corporation, or Adobe Systems, Inc.

References

1. S. Nasar, *A Beautiful Mind*, Touchstone/Simon and Schuster, New York, 1998.
2. R. Lougee-Heimer, "The Common Optimization INterface for Operations Research: Promoting Open-Source Software in the Operations Research Community," *IBM J. Res. & Dev.* **47**, No. 1, 57–66 (2003, this issue).
3. F. G. Gustavson, "Recursion Leads to Automatic Variable Blocking for Dense Linear-Algebra Algorithms," *IBM J. Res. & Dev.* **41**, No. 6, 737–755 (1977).
4. F. G. Gustavson, "High-Performance Linear Algebra Algorithms Using New Generalized Data Structures for Matrices," *IBM J. Res. & Dev.* **47**, No. 1, 31–55 (2003, this issue).
5. C. V. Apte, S. J. Hong, R. Natarajan, E. P. D. Pednault, F. A. Tipu, and S. M. Weiss, "Data-Intensive Analytics for Predictive Modeling," *IBM J. Res. & Dev.* **47**, No. 1, 17–23 (2003, this issue).
6. R. Agrawal, H. Mannila, R. Srikant, H. Toivonen, and A. I. Verkamo, "Fast Discovery of Association Rules," *Advances in Knowledge Discovery and Data Mining*, U. M. Fayyad, G. Piatetsky-Shapiro, P. Smyth, and R. Uthurusamy, Eds., AAAI/MIT Press, Cambridge, MA, 1996, pp. 307–328.
7. S. J. Hong, "R-MINI: An Iterative Approach for Generating Minimal Rules from Examples," *IEEE Trans. Knowledge & Data Eng.* **9**, No. 5, 709–717 (1997).
8. C. Apte, E. Bibelnicks, R. Natarajan, E. P. D. Pednault, F. Tipu, D. Campbell, and B. Nelson, "Segmentation-Based Modeling for Advanced Targeted Marketing," *Proceedings of the Seventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (SIGKDD)*, 2001, pp. 408–413.
9. C. V. Apte, R. Natarajan, E. P. D. Pednault, and F. A. Tipu, "A Probabilistic Estimation Framework for Predictive Modeling Analytics," *IBM Syst. J.* **41**, No. 3, 438–448 (2002).
10. C. Apte, E. Grossman, E. P. D. Pednault, B. Rosen, F. Tipu, and B. White, "Probabilistic Estimation Based Data Mining for Discovering Insurance Risks," *IEEE Trans. Intell. Syst.* **14**, No. 6, 49–58 (1999).
11. R. Anbil, J. J. Forrest, and W. R. Pulleyblank, "Column Generation and the Airline Crew Pairing Problem," *Documenta Mathematica, Journal der Deutschen Mathematiker Vereinigung, Extra Volume ICM III* (Invited Lecture at International Congress of Mathematicians, Berlin, 1998), pp. 677–686.
12. A. Bar-Noy, R. Bar-Yehuda, A. Freund, J. Naor, and B. Schieber, "A Unified Approach to Approximating Resource Allocation and Scheduling," *J. ACM* **48**, No. 5, 1069–1090 (2001).
13. A. Bar-Noy, S. Guha, Y. Katz, J. Naor, B. Schieber, and H. Shachnai, "Throughput Maximization of Real-Time Scheduling with Batching," *Proceedings of the 13th ACM-SIAM Symposium on Discrete Algorithms*, 2002, pp. 742–751.
14. M. X. Goemans and D. P. Williamson, "0.878-Approximation Algorithm for MAX-CUT and MAX-2SAT," *Proceedings of the 26th Annual ACM Symposium on the Theory of Computing*, 1994, pp. 422–431.
15. M. Shub, Editorial, *Foundations of Computational Mathematics* **1**, No. 1, 1–2 (2001).
16. L. R. Bahl, F. Jelinek, and R. L. Mercer, "A Maximum Likelihood Approach to Continuous Speech Recognition," *IEEE Trans. Pattern Anal. & Machine Intell.* **PAMI-5**, 179–190 (1983).
17. F. Jelinek, *Statistical Methods for Speech Recognition*, The MIT Press, Cambridge, MA, 1997.
18. D. E. Johnson, F. J. Oles, T. Zhang, and T. Goetz, "A Decision-Tree-Based Symbolic Rule Induction System for Text Categorization," *IBM Syst. J.* **41**, No. 3, 428–437 (2002).
19. M. Martens and T. J. Nowicki, "Ergodic Theory of One-dimensional Dynamics," *IBM J. Res. & Dev.* **47**, No. 1, 67–76 (2003, this issue).

20. C. Pugh and M. Shub, "Stable Ergodicity and Julienne Quasi-Conformality," *J. European Math. Soc.* **2**, 1–52 (2000).
21. K. Burns, C. Pugh, M. Shub, and A. Wilkinson, "Recent Results About Stable Ergodicity," *Proceedings of Symposia in Pure Mathematics, Vol. 69—Smooth Ergodic Theory and Its Applications*, A. Katok, R. de la Llave, Y. Pesin, and H. Weiss, Eds., American Mathematical Society, Providence, RI, 2001, pp. 327–366.
22. R. L. Adler, B. P. Kitchens, M. Martens, C. P. Tresser, and C. W. Wu, "The Mathematics of Halftoning," *IBM J. Res. & Dev.* **47**, No. 1, 5–15 (2003, this issue).
23. P. L. Bernstein, *Against the Gods, the Remarkable Story of Risk*, John Wiley & Sons, Inc., New York, 1996.
24. M. Crouhy, D. Galai, and R. Mark, *Risk Management*, McGraw-Hill Book Co., Inc., New York, 2001.
25. X. Guo, "Information and Option Pricing," *J. Quantitative Finance* **1**, No. 1, 38–44 (2001).
26. D. L. Jensen and A. J. King, "Frontier: A Graphical Interface for Portfolio Optimization in a Piecewise Linear-Quadratic Framework," *IBM Syst. J.* **31**, No. 1, 62–70 (1992).

Received June 24, 2002; accepted for publication November 11, 2002

William R. Pulleyblank *IBM Research Division, Thomas J. Watson Research Center, P.O. Box 218, Yorktown Heights, New York 10598 (wp@us.ibm.com).* Dr. Pulleyblank is the Director of Exploratory System Servers in the IBM Research Division and the Director of the IBM Deep Computing Institute. He has also served as the Research relationship executive responsible for the Financial Services sector in IBM, the Utility and Energy Services industry, and for the Business Intelligence group. Before joining IBM Research in 1990, Dr. Pulleyblank was the holder of the Canadian Pacific Rail/NSERC Chair of Optimization and Computer Applications at the University of Waterloo. He is a member of a number of boards, including the Mathematical Sciences Board of the NRC, the External Advisory Board of DIMACS, the iCORE Board of Directors, the Advisory Council of the Pacific Institute for the Mathematical Sciences (PIMS), and a member of the Scientific Advisory Panel of The Fields Institute for Research in Mathematical Sciences. In addition he serves on the editorial boards of a number of journals. Dr. Pulleyblank's personal research interests are in operations research, combinatorial optimization, and applications of optimization. In addition to writing a number of scientific papers and books, he has consulted for several companies, including Mobil Oil (helicopter routing); Marks and Spencer (depot management); Statistics Canada (survey validation); and CP Rail (train scheduling).