

SC08 参加報告

防衛大学校 応用物理学科 萩田 克美

2008 年 11 月 18 日～20 日、アメリカ合衆国テキサス州オースティンで開催された SC08 に参加した。SC は、1988 年以来、毎年開催されている「世界最大のスパコンに関する会議・展示会」である。SC は、ベンダーを中心とした展示会がメインであり、スパコンをユーザーとして利用する私にとっては、まさに、自分の買い物とは関係ない先端的な車や個別の部品が展示される「東京モーターショー」にでも来ている感覚である。展示会と並行して、HPC に関連する先端的な研究を報告する国際会議も催されている。ここでは、理研のレポーターとして、いくつかの研究報告を聴講した結果を中心に報告する。本報告では、HPC に興味があり多様な専門を背景にもつ方々が読者であると想定し、記録的価値を高めるとともに SC08 の雰囲気伝えるために、写真を多く用いることとした。偶然ではあるが、私は、昨年度の SC07 と欧州で毎年開催されている ISC08 にも参加している。それらの経験を踏まえつつ、本報告を読んでいただける方にとって参考になる情報提供ができれば幸いである。

昨年から注目されていた D. E. Shaw リサーチの Anton の動向であったが、事前のアブストラクト Web 上では見当たらなかった。D.E. Shaw の関連した研究の発表によれば、2008 年 10 月に 512 ノードでのプロトタイプが動作したとのことであり、アナウンスは控えめであった。25,000 原子のたんぱく質の MD についてミリ秒の計算を実現するのは、まもなくといったところのようである。

SC では、国際会議に先立ち、HPC 教育セッションとチュートリアルが開催されている。本年度は、参加することができなかったが、価値の高い催しなので、1 年さかのぼって簡単にレポートしておく。SC07 の教育セッション (Educational Program) では、全米における HPC 関連の初等中等教育についての発表がなされると共に、スパコン利用を支援する Web ベースアプリなどのハンズオンでの講習会が催されていた。また、SC07 のチュートリアルでは、CUDA プログラム、MPI と OpenMP のハイブリッドプログラム、OpenMP 3.0 など、最新のトピックが、実例と経験を基に、日本国内の講習会などでは聞けないハイレベルな内容があり、実際にスパコンを使う研究者・技術者にはお勧めである。

さて、今年の SC08 の報告に移ろう。

Registration は、昨年度同様、PC で各自参加登録し、有人ブースでネームプレートを受け取るものである。ネームプレートには、磁気テープ付いており、展示ブースに設置された読み取り機を通じ個人情報が発示者に引き渡される仕組みとなっている。このあたりのシステムティックさは、1 万人の参加者を裁く「The 展示会！」といった感じである。展示会の顔ぶれは、昨年と大体同じである。日本からの参加としては、地球シミュレータを ES2(NEC SX-9)にリプレイスする JAMSTEC が復活参加している以外、TOP500 に絡むようなスパコンを有する大手の HPC ファシリティが継続的に参加している。ロスアラモス研究所、NASA、オークリッジ研究所、アルゴンヌ研究所などのペタ flops 級のスパコンが導入されている大型ファシリティのブースは、非常に大きく、目立っていた。相対的に理研のブースが小さく感じたのが率直な印象である。

初日、8:30 より、Welcome セッションと Keynote セッションが開始された。SC20 周年の会議ということもあり、Gordon Bell を始め、スパコン業界の著名人が、これまでのスーパーコンピュータの歴史語る 10 分間のビデオが上映された。その中で、日本の計算機としては地球シミュレータ誕生の衝撃とその成果について絶賛されていた。あらためて地球シミュレータのスパコン史における位置づけの重要性を認識した。最初の Keynote としては、Dell の創始者 Michael Dell が、Dell の HPC への Contribution について講演した。ベクトル、マルチプロセッサと HPC が進化し、現在では、x86 などの標準化された HPC が、TOP500 に占める割合を高めるに至り、Dell が活躍しているという骨子である。

その後、10 時半からは、大注目のロスアラモス研究所の Roadrunner の講演であった。Roadrunner のシステム詳細が紹介され、各アプリケーションでの性能評価の結果が示された（詳細は、聴講報告①参照）。アーキテクチャの説明から受けた印象では、PlayStation3 の MPI プログラミングの延長として対応できるとの感触を得た。次の講演では、NASA が所有する最新計算システムのアプリケーション毎の性能比較についての紹介であった（詳細は、聴講報告②参照）。CPU の単体性能としては Xeon よりも Power5+や Itanium2(Montvale)がよいという当然の結論であるが、SGI ICE では CPU 数を増やして性能を出している。Roadrunner と SGI ICE の両方とも、LINPACK に対する経済性は十分であるが、アプリケーション開発に対する経済性について苦しんでいるのではないかという感覚を覚えた。

午後は、「HPC in Transportation」セッションに参加した。Goodyear と Boeing の発表であった。両者とも、HPC を駆使して開発をしている現状を事例紹介した。産業利用のためには、実験とシミュレーションを確実に連携させている点は学ぶところが大きいと考えた。また、Goodyear は、エネルギー省の CRADA

プログラムの下、Sandia 研究所と長年共同研究を進め、着実に成果をあげている点も注目すべきであると感じた（詳細は、聴講報告③参照）。

休憩を挟み、「HPC in Finance」セッションに参加した。Morgan Stanley のみの発表であった。当初は、Golden Tree Asset Management も発表する予定であったが、キャンセルになっていた。一連の金融危機が関係あるのかもしれませんが。Morgan Stanley の発表では、リスク計算のモンテカルロ法などでの入力条件に対する分岐により計算が膨大になるので、CPU の数が勝負となる旨強調していた。

1 日目の夕方のイベントは、最新の TOP500 リストの発表とポスター発表である。TOP500 の BOF で配布された「The TOP500 Report」に記載される TOP25 から、ついに、日本勢が姿を消している。復活することを期待するのみである。ちなみに、TOP9 まで USA、10 位が中国である。1 位、2 位が 1 ペタ級、3 位から 6 位が 0.5 ペタ級、7 位～11 位（ドイツ）が、0.2 ペタ級である。今回の注目は、2008 年 6 月 (ISC08) で初めて 1 ペタ flops 超えて 1 位となった Roadrunner と前評判の高い Cray XT5 の勝負であった。結果として、ISC08 からさらにアップグレードした Roadrunner が僅差で死守した形である。

TOP500 の発表の後は、ポスター発表に参加した。ポスター発表は、あらかじめ注目していた 3 件を中心に回った（詳細は、聴講報告④～⑥参照）。プログラムのチューニング、GPGPU や CELL のプログラミング、CPU やネットワークシステムの性能評価など、HPC の基盤技術に関する基礎的な検討に関する発表が幅広く行われている。正直に言えば、大半の発表が専門的すぎ理解することができないものであった。このような専門的で地道な研究が結集して、HPC ができているということを改めて感じた。

2 日目、午前は、癌に対するテーラーメイド医療のための情報供給システムの構築の話と、UC バークレーにおける新たな並列計算に関する研究センターについての紹介があった。前者については、巨大な情報システムの構築という点では理解できるが、大規模計算という意味での HPC がどこに組み込まれてくるのか良く分からなかった（詳細は、聴講報告⑦参照）。Grid で情報収集し分析を通じて業務に活かしていくことはあらゆる分野で成り立つ手法であり、先駆的な取り組みで一般の技術レベルが向上していくことは価値が高いと考えられる。後者の講演では、100Core 以上の並列ソフトウェアの作成を目標に、84 台の FPGA を用いた 1008Core の RISC エミュレーション環境を構築して、新しい並列計算に関する研究を進めているとのことである（詳細は、聴講報告⑧参照）。

その後、Blue Gene の BOF に参加した（詳細は、聴講報告⑨参照）。実質的には、エネルギー省の INCITE プログラムで実施された計算事例紹介であった。

午後、Gordon Bell 賞の Finalists の発表を聴講した。今年度の Gordon Bell の Finalists は、すべて、アメリカの最新スパコンを利用し、数万の CPU を利用した並列計算である。昨年は、日本からも地球シミュレータや MD-GRAPE3 などで Finalist を排出していたが、今年はずいぶんアメリカ以外は手を出せない状況といった感じである。最初の講演は、地震波の進行に関するシミュレーションを Texas の Ranger システムで 62,976 個の Core を用いて並列計算したものであった。論文の PDF は、<http://scyourway.nacse.org/conference/view/gb109> を通じ、IEEE のアーカイブか ACM のデジタル図書館から入手可能である（以下同じ）。次の講演は、オークリッジ研究所の高温超伝導の不純物効果に関するシミュレーションであった。この研究では、Cray XT4 の 49,282CPU（2/3 がオークリッジ研究所のスパコン、1/3 はテネシー大学のスパコン）を用いて、ハバード模型の量子モンテカルロ計算を行い、409TFlops の性能を示したものである。論文は、<http://scyourway.nacse.org/conference/view/gb118>。さらに、Cray XT5 と Cray XT4 を同時利用し、1.3 ペタ Flops の実効性能を出したとのことである。本日最後の講演は、Texas の Ranger を用いたマントル対流の計算であった。論文は、<http://scyourway.nacse.org/conference/view/gb132>。32,786Core の並列性能はよいが、62,464Core ではパフォーマンスが落ちていた。

夕方は、SGI ICE ユーザーの BOF に参加した（詳細は、聴講報告⑩参照）。1Core あたりのメモリが小さく苦労しているというのが全体を通じての印象であった。

最終日、冒頭で、次回 SC09 が開催されるポートランドの説明があった。富士山のようなきれいな山があるということが印象に残った。さて、午前は、「エネルギーと HPC」というテーマで2つの講演を聴講した。最初の講演では、エネルギー問題全般の話と、ITER、燃焼、材料などの分野で HPC が重要な役割を果たすであろうことを力説していた。次の講演では、CO₂ 貯蔵技術におけるシミュレーション事例についての紹介がなされた。両講演から、エネルギー問題において HPC の果たす役割は大きいと感じた。

その後、Mathematica の MPI 拡張に関する Exhibiter Events に参加した（詳細は、聴講報告⑪参照）。

引きつづき、Gordon Bell 賞の Finalist の講演を聴講した。本日最初の講演は、Roadrunner を用いたレーザープラズマのシミュレーションであった。論文は、<http://scyourway.nacse.org/conference/view/gb121>。2 番目の講演は、引き続き Roadrunner を用いた MD シミュレーションであった（詳細は、聴講報告⑫参照）。論文は、<http://scyourway.nacse.org/conference/view/gb124>。Cell 上での MD プログラムのチューニングをしっかりとやっている印象を受けた。また、個

人的には、Cell-SPE と Opteron の間の通信が多いような気がしたが、実際にやってみなと判断がつかないと思われる。最後の講演では、ローレンス・バークレー研究所の Cray XT5 などを用いた LS3DF というフラグメント法の電子状態計算であった。論文は、<http://scyourway.nacse.org/conference/view/gb115>。この計算では、Cray XT4 の 36,864Core 並列で 135Tflops、Blue Gene/P の 163,840Core 並列で 224Tflops、Cray XT5 Jaguar の 147,456Core 並列で 442Tflops の性能を示すとのことである。その後、HPC Advisory Council Initiative の BOF に参加したが、収穫はなかった（詳細は、聴講報告⑬参照）。

午後は、SC Award の発表であった。今年度の Gordon Bell 賞（ピーク性能）は、オークリッジ研究所の「New algorithm to enable 400+ TFlop/s sustained performance in simulations of disorder effect in high-Tc superconductors」であった。この研究では、Cray XT4 の 49,282Core を用いて、ハバード模型の量子モンテカルロ計算を行い、409TFlops の性能を示したものである。Special Award (Algorithmic Innovation) は、ローレンス・バークレー研究所の発表「Linear Scaling Divided-and-conquer Electronic Structure Calculations for Thousand Atom Nanostuctures」が受賞した。

午後の後半は、「Applications: Models and Analysis」というセッションに参加した。このセッションでは、地震のシミュレーション、MD 軌道の並列解析、CELL を用いた並列推定計算の 3 件の講演があった。MD 軌道の並列解析については、D. E. Shaw リサーチの発表であり、Anton の動向が得られるであろうと期待していた。実際、Anton 関連では、プロトタイプが動作したとの情報のみであった。講演の内容は、分散入力の解析ツールを整備したとのことであった（詳細は、聴講報告⑭参照）。

最終日の夕方、会議のディナーがテキサスの牧場で開かれた。バスに乗るまで約 1 時間、そして 30 分の移動後についた牧場は、難民キャンプの様相であった。とにかく、この時期のテキサスは寒い。明らかに 10℃よりも低い。まずはお酒を確保し、列に並ぶこと 30 分、やっと食料にたどりつくという感じであった。井上室長、野村技術参事、南チームリーダーには大変お世話になりました。

最後に、SC08 への参加の機会をいただいた理化学研究所次世代スーパーコンピュータ開発実施本部にお礼申し上げます。特に、出張の支援をしていただいた内田さんにはお世話になりました。また、本出張のために、授業の代理をしていただいた防衛大学校応用物理学科の同僚には、感謝の意は絶えません。そして、同じ理研レポーターの立場で参加された斉藤さん、牧野さんには、お世話になりました。特に、旅程を同じくした斉藤さんには、多くの楽しい時間とドキドキした時間を共有させていただきありがとうございました。

●聽講報告

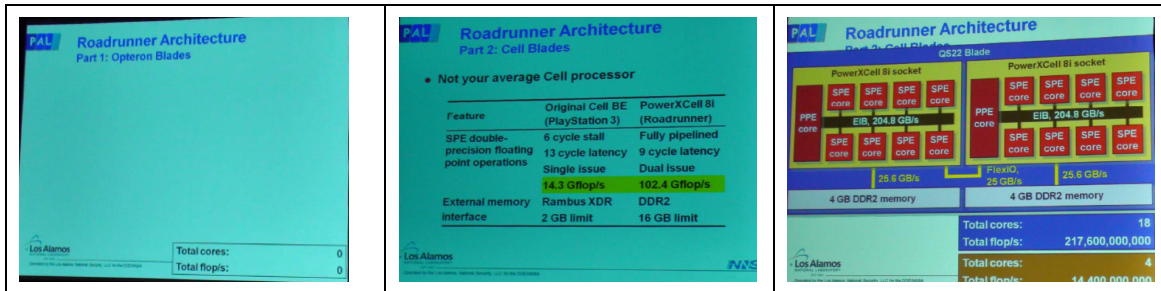
① Papers 「Entering the Petaflop Era: The Architecture and Performance of Roadrunner」 by Kevin Barker et. al. (Los Alamos National Laboratory)

ロスアラモス研究所に導入され、今回の TOP500 で 1 位となった Roadrunner についての紹介であった。アーキテクチャの説明からは、PlayStation3 の MPI ベースのプログラミングスキルがあれば、対応できると感じた。Opteron と Cell の両方を活用してアプリケーション性能を引き出すことに苦労している様子が伺えた。Opteron と Cell の PPE との通信性能がネックになっていると感じた。また、3.6Gflops の Opteron と 102.4Gflops の PowerXCell がアプリレベルの性能が同じオーダーというのは、アプリが Opteron よりに造られているためであり、PPE-Opteron の通信律速に加え、アルゴリズム上の課題も大きいと感じた。

論文の PDF は、<http://scyourway.nacse.org/conference/view/pap333>。

発表のポイントは次の通りである。

- ・ Roadrunner とは、ロスアラモス研究所に設置された 1.38Pflops@peak、1.026Pflops@LINPACK の性能を持つ計算機である。12240Core の Opteron と 12240Core の CELL のハイブリッド計算機。ピーク性能が倍精度で 1.38 ペタ flops、単精度 2.9 ペタ flops。1Core のメモリは 4GB で総メモリは 98TB。(実際の計算には、28TB が Opteron、52TB が Cell に割り当てられている。)
- ・ PPE と Opteron の間の通信が遅くてもよいとの考えでシステム構築を行っている。これは意外にも思えるが、Playstation3 での MPI 通信を扱うプログラムを考え、Opteron レイヤーで汎用的な計算を行うとすれば、CELL の性能を引き出すために必須の SPE と PPE の通信と Opteron 間の通信が速ければよいということは納得である。
- ・ しかしながら、現行システムでは、PPE と Opteron の通信がボトルネックである。PPE と Opteron の通信が、Opteron 間の通信に比べ、相対的に性能が上回るあたりで、よい性能を示している。
- ・ 7 種類のアプリケーションを Roadrunner の初期アプリと設定し、その評価結果を発表した。CELL を使えば、Opteron のみに比べ 2~6 倍速くなる。



Roadrunner Architecture Part 1: Scaling Out

Total cores: 480
Total flops: 5,395,200,000,000

Roadrunner Architecture Part 4: Scaling Out

Total cores: 7,200
Total flops: 80,928,000,000,000

Roadrunner Architecture Part 4: Scaling Out

Total cores: 122,400
Total flops: 1,375,776,000,000,000

Roadrunner Architecture Summary of Key Characteristics

- Hybrid architecture**
 - 12,240 Opteron cores for control- or network-intensive routines and irregular memory accesses
 - 12,240 Cell sockets (97,920 SPE cores) for compute-intensive routines with regular memory accesses
 - Equal memory (4 GB) per Opteron core and Cell socket
 - Total of 96 TB memory
- High peak performance**
 - DP peak: 1.38 Pflop/s
 - SP peak: 2.91 Pflop/s
- Ordinary InfiniBand network**
 - Approximately a fat tree—see paper for details
- Modest # of nodes (3,060)**
 - 91% of performance comes from SPE cores

Why This Architecture?

- Attempt to optimize application performance given multiple constraints**
 - Cost
 - Flexibility
 - Power & cooling
 - Floor space
 - Delivery schedule
- Roadrunner's hybrid architecture deemed the best solution**
- Hybrid may be the new trend in HPC**
 - 1970s: HPC is scalar; LANL is an early adopter of vector (Cray 1 #1)
 - 1980s: HPC is vector; LANL is an early adopter of data-parallel (TMC CM-1)
 - 1990s: HPC is data-parallel; LANL is an early adopter of distributed-memory (TMC CM-5)
 - 2000s: HPC is distributed-memory; LANL is an early adopter of hybrid (Roadrunner)

Memory Subsystem Performance

- PPE core provides 78% more peak flops than Opteron core — but —
- PPE observes only 16% of Opteron's memory bandwidth
- Our experience**
 - PPE not fast enough for significant computation
 - PPE speed on small kernels is about 1/3 Opteron speed
 - On Roadrunner, PPEs are best used for shuttling data between SPEs and Optérons

Communication Subsystem Performance

- Key contributor to performance of many parallel applications**
- Complexities of measuring on Roadrunner**
 - Multiple networks: Element Interconnect Bus, FlexIO, PCI Express, HyperTransport, InfiniBand
 - Different low-level protocols: put/get vs. send/receive
- Our approach**
 - Normalize all protocols to send/receive
 - MPI send/receive for Opteron-Opteron communication
 - DaC2 send/receive for PPE-Opteron communication
 - Cell Messaging Layer for SPE-SPE communication
 - Measure ping-pong performance (half round-trip time) across each interconnect type

Small-Message Communication Time

Large-Message Communication Time

- PPE-Opteron is bandwidth bottleneck
- Same link technology as Opteron-Opteron (PCI Express x8)
- As SW matures, we expect performance to improve to current MPI-B performance

Early Roadrunner Applications

- VPIC**
 - Particle-in-cell code
 - 7X improvement over Opteron-only
- SPaSM**
 - Short-range molecular dynamics code
 - 6X improvement over Opteron-only
- Milagro**
 - Implicit Monte Carlo code
 - 6X improvement over Opteron-only
- PetaVision**
 - Neuron & synapse simulation
 - 1.144 Pflop/s (single proc.)
- Each ported to Roadrunner by a couple of people in a short period of time**
 - Months, not years
 - Most had to learn Cell programming first
 - All had to deal with preproduction HW & SW
- Relatively few code changes**
 - 10–30% of code (rest.)
- Yes, Roadrunner is programmable**

Challenge: Sweep3D

- Neutron-transport kernel
- 3-D global grid with 2-D data decomposition
- Wavefront communication**
 - Receive boundaries from upstream
 - Compute
 - Send boundaries downstream
 - Repeat from each of eight corners
- Hard to get performance at scale**
 - Small messages (few KB)
 - Tightly coupled—runs only as fast as slowest link
 - Tradeoff between frequency of communication and available parallelism

Sweep3D on Roadrunner

- All compute done on SPEs**
- PPEs and Optérons used as smart NICs**
 - (Remember: 91% of performance on SPEs)
- Same basic data structures and control flow as conventional Sweep3D**
 - Cell Messaging Layer provides MPI for SPEs
 - One MPI rank per SPE
 - Test Roadrunner as a 97,920-SPE cluster
- 2X improvement at scale**

Conclusions

- Roadrunner is a complex architecture**
 - Three types of processor cores
 - Multiple internal interconnects
 - Many address spaces per node
- Good performance is possible, however**
 - Even communication-heavy Sweep3D sees 2X improvement over Opteron-only runs with 4X expected in the future
 - Other applications see 6–7X improvement
- Hybrid computing may be here to stay**
 - Higher performance than other systems
 - Scales well (few powerful nodes enables smaller networks)
 - Opens up exciting new research possibilities in system architecture, programming models, performance tools, and more

Additional Roadrunner Resources

- Read the paper
- Visit our Web site: <http://www.lanl.gov/roadrunner>
- Download our software: <http://cellmessaging.sf.net/> (Cell Messaging Layer)
- Stop by our booths
 - Los Alamos National Laboratory (#550)
 - NNSA Advanced Simulation and Computing (#512)
- Attend our talks
 - Birds of a Feather session (Wednesday evening)
 - "Roadrunner: First to a Petaflop, First of a New Breed"
 - ACM Gordon Bell Finalists session (Thursday morning)
 - "1.374 Pflop/s Trillion-particle Particle-in-cell Modeling of Laser Plasma Interactions on Roadrunner"
 - "369 Tflops Molecular Dynamics Simulations on the Roadrunner General-purpose Heterogeneous Supercomputer"

② Papers 「Scientific Application-based Performance of SGI Altix 4700, IBM POWER 5+, and SGI Altix ICE 8200 Supercomputers.」 by Subhash Saini et. al. (NASA Ames Research Center)

NASA に導入された SGI ICE のアプリケーション性能について、POWER5+ や Altix(Itanium 2)と比較結果の発表である。NASA の SGI ICE 8200EX は、今回の TOP500 の 3 位に入賞した、約 0.5 ペタクラスの計算機である。本検討は、このシステムの一部を利用したものと考えられる。同じ数の Core を用いた計算では、POWER5+が最も性能がよく、SGI ICE は、半分程度の性能が出るというものであった。NASA の SGI ICE は、LINPACK 性能は良いものの、1Core あたりのメモリが 1GB であるなど、アプリケーション実行のためには使い勝手が良いとはいえないようである。

論文の PDF は、<http://scyourway.nacse.org/conference/view/pap452>。

発表のポイントは次の通りである。

- NASA が所有する 3 種のスパコン (SGI Altix 4700、SGI ICE 8200、Power5+) について、アプリケーションベースで評価した。
- SGI ICE は、低価格な Xeon を使い、Core 数を増やして LINPACK 性能を稼いでいる機種である。NASA に導入された SGI ICE の 1Core あたりのメモリは 1GB である。Core のネットワークポロジは 4 次元ハイパーキュービックであり、 $128\text{core} \times 16 = 2048\text{core}$ が 1 つの unit となっている。
- Core の単体性能としては Xeon よりも Power5+や Itanium2(Montvale)がよいという当然の結論である。特に、メモリバンド幅の差が顕著である。
- NASA が利用している 4 つのアプリケーションについて性能評価した結果からは、アプリケーション性能は Power5+が一番である。SGI ICE はメモリバンド幅が小さいために、アプリケーション性能が出ず、Power5+や Altix の半分程度の性能である。なお、512Core までの評価である。
- また、MPI と OpenMP のハイブリッドコードは、Altix や ICE では性能が出せないようである。これらでは、ハイブリッドでは性能が出せず、MPI だけのコードにした方がよいようである。

Outline

- HEC systems
 - SGI Altix 4700 density (dual-core)
 - SGI ICE 8200 (quad-core)
 - IBM POWER5+ (dual-core)
- HPC challenge benchmarks (HPCB)
- Kernels & compact applications (NPB)
- Applications
 - Computational fluid dynamics
 - OVERFLOW-2 (structured mesh)
 - CART3D (structured mesh)
 - USM3D (unstructured mesh)
 - Climate
 - ECCO
- Results [MPI & Hybrid (MPI+ OpenMP)]
- Conclusions

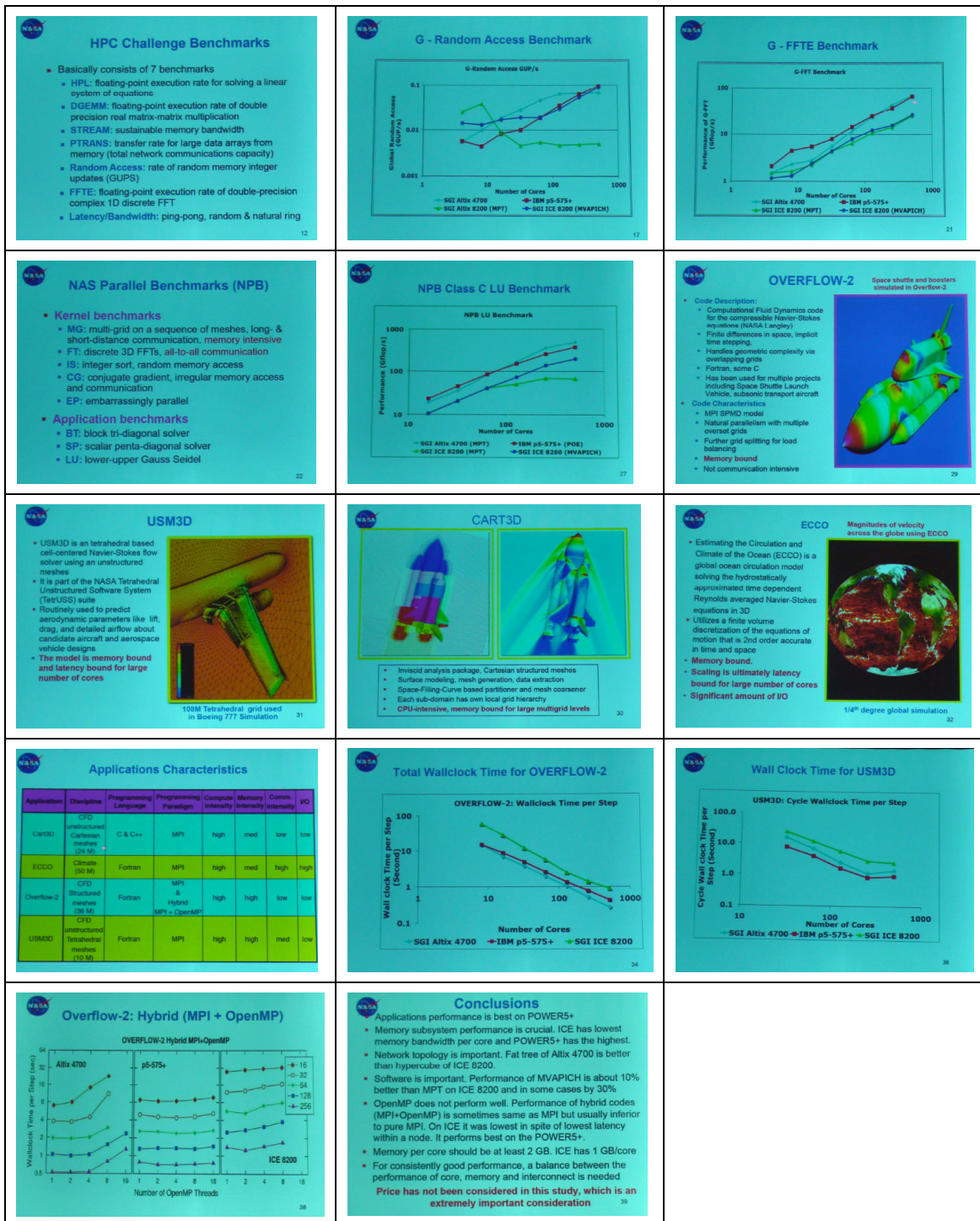
SGI Altix 4700, ICE 8200 and POWER5+

Characteristics	Altix 4700 Montvale	ICE 8200 Clovertown	POWER5+
Processor	Intel Xeon 2 Montvale	Intel Xeon X3355 Clovertown	IBM POWER5+
Cores per socket	2	4	2
Cores per node	4	8	16
Core clock (GHz)	1.67	2.66	1.5
Flap/clock/core	4	4	4
Peak per core (GFlops)	6.67	10.64	7.5
L1 cache (KB)	32	32 (I) & 32 (D)	64 (I) & 32 (D)
L2 cache	256 KB	4 MB shared by 2 cores	1.50 MB (HIO) shared
L3 cache	9 MB/core	none	36 MB (off-chip)
Local memory/core (GB)	2	1	2
Interconnect	NUMALink4	InfiniBand	HPS (Federation)
Interconnect topology	Fat tree	Hypercube	Multi-stage
Total number of cores	1,024	4,096	640

SGI ICE 8200 Network Topology

4D Hypercube

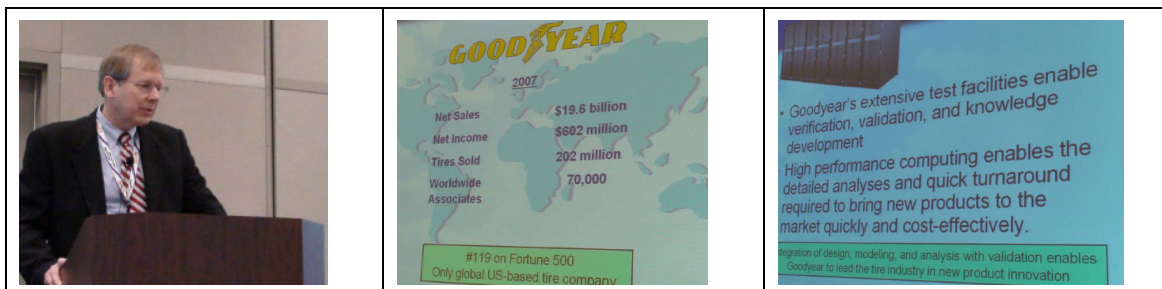
- Each vertex is an entire one IRU (16 compute nodes or 128 cores)
- One cube has 1024 cores (8 vertices x 128 cores)
- Two cubes has 2048 cores.

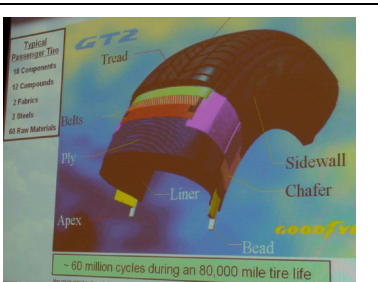
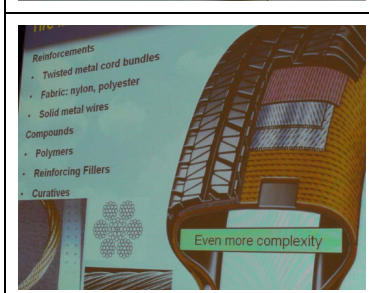
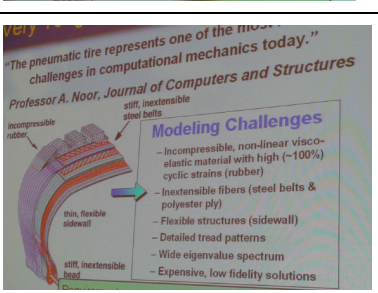
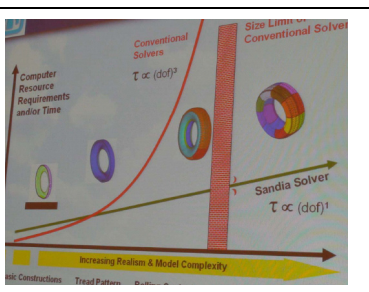
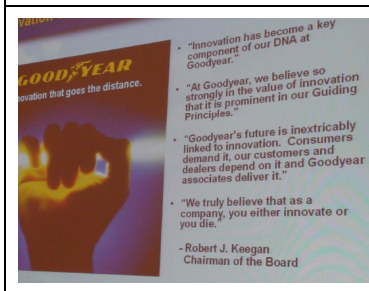


③ Masterworks 「Enabling Analysis-based Tire Development with High Performance Computing」 by Loren Miller (Goodyear Tire & Rubber Company)

Goodyear が、1993 年より継続しているエネルギー省の CRADA プロジェクトの下で、有限要素法(FEM)によるタイヤのトレッドパターンの大規模シミュレーション手法を、Sandia 研究所との共同研究として開発した経緯に関する発表である。企業単独では解決できなかった課題を、Sandia のコードを用いることで解決し、実際の製品開発に大きく役立っていることはすばらしいと感じた。発表のポイントは次の通りである。

- タイヤの開発において、燃費性能と安全性の向上が重要な課題である。
- Goodyear は、HPC を、詳細に解析する手段であるとともに、低コストで短期間に市場に製品を送り出す手段として捕らえている。
- タイヤのモデリングは、困難で複雑である。タイヤは、60 種類の材料から多数のコンポーネントから構成されている。特に、ゴム材料は、高分子や補強のためのフィラーなどからなり、Mullin's effect、Payne effect など複雑な性質を示し、扱うのが難しい。
- 1994 年に Goodyear で計算できる FEM モデルの自由度の最大値は 90,000 自由度であった。タイヤトレッドのモデルに必要な最小のサイズは 250,000 自由度と見積もられ、手が出る状況ではなかった。そこで、1993 年より、Sandia 研究所との共同研究で、エネルギー省の CRADA(Cooperative Research and Development Agreement)プロジェクトを実施した。Sandia 研究所が作成したソルバーを用いることで、計算コストを大幅に減らすことに成功し、大規模な FEM 計算が可能になった。
- タイヤトレッドの FEM のシミュレーションにより、タイヤ製品の開発期間を 3 年から 1 年に短縮することに成功した。
- 現在では、128 ノードの Linux クラスターを利用しており、RunFlat タイヤの開発で、非常に活躍している。



ENGINEERING - Predictive testing was the more "comfortable" approach. - Lab tests were developed to substitute for lengthier, less controllable road tests. - Corollary benefits: - Provided basis for the development of engineering design criteria in support of modeling - Provided data to validate models Modeling and analysis was less "comfortable". - However, it was predicted to become less expensive and more capable over time Why were many engineers skeptical?	Since 1898, we'd always developed tires experimentally. - Modeling tires is difficult and complex. - Existing models (1992) were too crude to make design decisions and took too long to run. - Knowledge and time were lost at the designer-to-analyst handoff. Just use a [commercial code] and use your PhD's for tech support... Analysis won't give us a competitive advantage anyway...	
TIRE - Reinforcements - Twisted metal cord bundles - Fabric: nylon, polyester - Solid metal wires - Compounds - Polymers - Reinforcing Fillers - Curatives 	Rubber Stress-Strain Curve (quasi-static) Pseudo-elastic behaviour - Stress softening on unloading vs. loading associated with virgin state (Mullin's effect) - Stress softening due to overstraining - Pre-conditioning (3 to 20 cycles) - Hysteretic stress-strain cycling - Residual strain Dynamic Mechanical Analysis Non-linear visco-elastic behaviour Dependence on: - Strain (Payne effect) - Temperature - Frequency - Strain history Rubber is the most complex component	Very difficult "The pneumatic tire represents one of the most challenging in computational mechanics today." Professor A. Noor, <i>Journal of Computers and Structures</i>  Modeling Challenges - Incompressible, non-linear visco-elastic material with high (~100%) cyclic strains (rubber) - Inextensible fibers (steel belts & polyester ply) - Flexible structures (sidewall) - Detailed tread patterns - Wide eigenvalue spectrum - Expensive, low fidelity solutions
SUMMARY - Model size - Largest model ever run at Goodyear in 1994: 90,000 degrees of freedom - Static, smooth, axisymmetric model ran for months - Solution times increased as cube of the model size - Estimated minimum model size required to model tire tread wear: 250,000 degrees of freedom - Solution time estimated at 15.6 years using more memory than CRAY ever put in the 1-MP - In contrast, build and test cycle for tread wear required four to six months - Tougher yet! Dynamics vs. statics - Not much happens sitting still in a parking lot! - Aircraft takeoff and landing - Taxing - Cornering, braking, accelerating - Pot holes Our goal: one week tread wear analysis	SUMMARY - No commercial solver (1994) was capable of solving this problem in the time required - Two alternatives for dynamics were axisymmetric (smooth tire) solutions only - Internal parallel code - Tire industry consortium through COMCO Too slow and, if we didn't need tread patterns to achieve tire performance, we wouldn't put them on! Now what?	
SUMMARY - One of the key topics of discussion between Goodyear and Sandia at the beginning of the relationship was what benefits each organization would receive. - Goodyear's benefits were illustrated in the preceding slides. - Three years after the CRADA began, Al Narath, then laboratory director, testified before the Congress that, as a result of the relationship, Sandia was able to solve previously intractable weapons design problems. Good for Goodyear Good for national security		designer/analyst - Designer to analyst handoff was the way finite element analysis was done. - Since the analyst begins only after the design is "complete," the process is inherently serial. - Design intent is never fully communicated to analyst. - Designer forced to put his thoughts on hold until analyst completes lengthy process of model making and computer runs Solution: Goodyear designed and developed its own, automated design/analysis system - Standard defaults for specific types of analysis - Over-rides available for analysts - Creates model, applies boundary conditions, submits job to Linux clusters, post-processes results
GOODYEAR Innovation that goes the distance.  "Innovation has become a key component of our DNA at Goodyear." "At Goodyear, we believe so strongly in the value of innovation that it is prominent in our Guiding Principles." "Goodyear's future is inextricably linked to innovation. Consumers demand it, our customers and dealers depend on it and Goodyear associates deliver it." "We truly believe that as a company, you either innovate or you die." - Robert J. Keegan Chairman of the Board		

- ④ Posters 「Performance of Electronic Structure Calculations on Blue Gene and Cray XT4 Computers」 by Andrey Asadchev et. al. (Iowa State Univ. & Ames Lab.)

GAMESS の中にある 2 次の摂動で電子相関を簡便に見積もる方法である MP2 法のチューニングについての発表である。このチューニングでは、BLAS2 と BLAS3 を駆使して行列変換を行うことがポイントである。20 アミノ酸からなる TRP タンパクについて、425 の active 軌道と 1107 の仮想軌道を用い計算した結果である。変換行列に 1.25TB のメモリが必要である。結果的には、128Core の Blue Gene で、64Core の Cray XT4 と同じ程度の計算時間で計算可能になったとのことであった。まだ改良の余地があり論文未発表とのことである。発表者に、Cray と Blue Gene のどちらの方がアプリケーション性能を出せるか質問したところ、当然 Cray であるという回答であった。

- ⑤ Posters 「Caching on a Chip Multi Vector Processor」 by Akihiro Musa et. al. (東北大学 & NEC)

現状のベクトル CPU はマルチコア化されていない上、キャッシュが無いので、メモリバンド幅が勝負である。SX-10 あたりでは、マルチコア化は必須である。SX-6 相当の性能のベクトル CPU 4 つとメインメモリの間に Shared Vector Cache をつけてマルチコア化した試作品を作成し、東北大学の既存のマルチスレッド最適プログラム（ベクトル向き）を用い、メモリバンド幅の重要性を検証した。スケーラビリティは、バンド幅に依存していることが明確になった。Shared Vector Cache に関しては、Non-blocking cache と 32 sub-cache の機能だけでは不十分で、現状の SX-9 のベクトルキャッシュには実装されていない Miss Status Handling Register がないと、キャッシュヒット率が上がらないとのことであった。

⑥ Posters 「Storage Performance Analysis in ES and ES2」 by Ken'ichi Itakura et. al. (JAMSTEC)

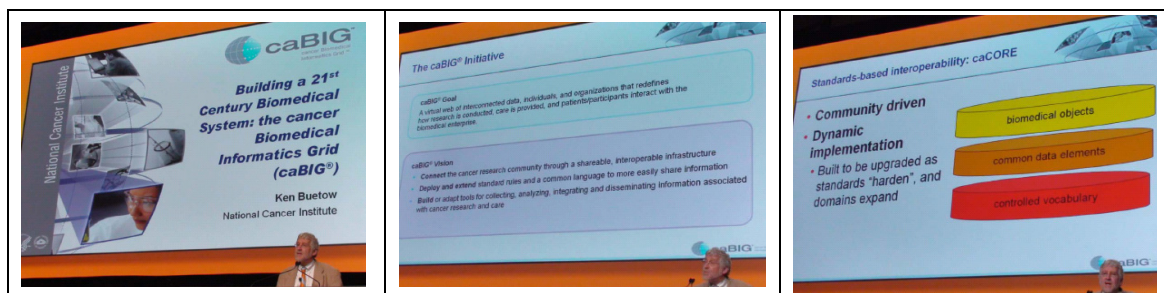
次期地球シミュレータ(ES2)のディスクシステムにおいて、Stage IN と Stage OUT の制御について検討した研究である。地球シミュレータでは、計算時の I/O 性能を高めるために、バッチジョブの実行前にローカルディスクにファイル転送 (Stage IN) し、バッチ終了後に出力ファイルを転送(Stage OUT)する機能を用いている。640 ノードの ES では、160 ノードを 1Gbps で 1 台の交換機に接続し、その交換機 4 台を 2Gbps でディスクシステムに接続している。160 ノードからなる ES2 では、16 ノードを 4Gbps で 1 台の交換機に接続し、その交換機 10 台を 4Gbps でディスクシステムに接続している。この研究では、ある 1 日の実際に地球シミュレータで流れたジョブのデータサイズと計算時間を用いて、Stage IN と Stage OUT にかかる時間について検討している。ES の運用データを分析すると、Stage OUT の時間については、ばらつきがあり、極端に遅いものがあった。原因は、Stage IN が優先的に帯域を使うためである。計算時間が半分で、データの量が同じと仮定する場合、Stage IN の時間は 37418sec から 9076sec に、Stage OUT の時間は 782sec から 94sec に削減することが分かったとのことである。データ量の増加が想定されること、Stage IN の優先度設定についての最適値の検討などが今後の課題であり、3 か月分のデータを用いてさらに検討を深めるとのことであった。

⑦ Invited Speakers 「Developing an interoperable IT Framework to Enable Personalized Medicine」 by Ken Bertow (National Cancer Institute)

癌患者に関する医療情報や解析ツールをがんセンターの間で共有する caBIG という grid システムについての紹介が行われた。実際にどのような成果が上がっているかの統計については紹介が無く、どこまで実用に至っているのかの情報を得ることはできなかった。

発表のポイントは次の通りである。

- 21 世紀の創薬としては、遺伝情報などに基づくテーラーメイド創薬を目指す考えである。
- バイオメディカル分野においては、情報が孤立していることが問題であった。臨床／画像診断／病理／分子バイオロジーの 4 分野の連携を促すために、caBIG という情報交換の枠組みとして、各分野の情報の API を整理して Grid ベースの Web サービスを構築した。
- caBIG は、56 の NCI のがんセンターと 16 の地域のがんセンターをつないでいる。National Health Information Network との接続や、英国、インド、中国、ラテンアメリカとの連携も予定している。
- caGrid、NHLBI Grid、NCTI ONIX の 3 つのグリッド情報は、相互に情報交換を行い、Grid のグリッドを実現している。
- ユーザーに対しては、single web のポータルサイトを構築している。Gene 情報、Genome 情報、臨床情報および、解析ツールを提供している。
- テーラーメイド創薬のために、BIG Health Consortium を設立した。政府、企業、個人を結んでいる。



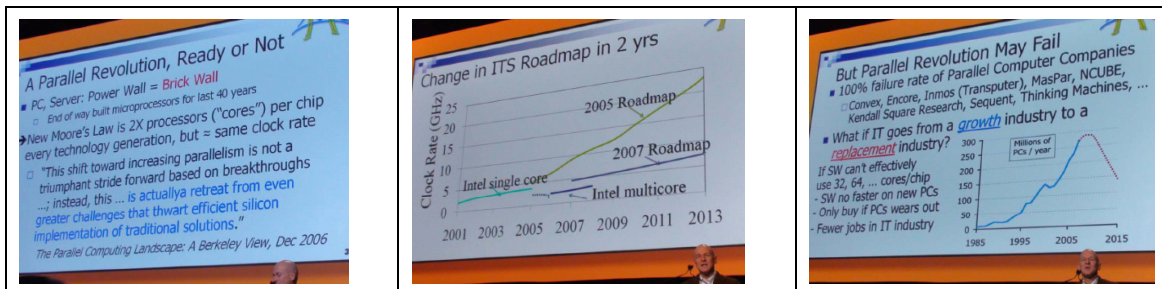


⑧ Invited Speakers 「Parallel Computing Landscape: A View from Berkley」
by Dave Patterson (U.C. Berkeley)

UC バークレーに設立された新たな並列ソフトウェアの研究センターの紹介であった。CPU の単体性能の成長は止まり、CPU 数の増加で性能を上げていることから、FPGA を用いて 90MHz の RICS チップ 1008Core の並列環境を構築し、並列ソフトウェアの研究を行っているとのことである。

発表のポイントは次の通りである。

- CPU や半導体のロードマップは、2005 年に比べ、2007 年では大幅に下方修正されている。キラーCPU の登場は期待することはできない。プログラムを早くするには、並列ハードウェアの利用が必須である。今後、IT 産業は成長産業から、置き換え産業になると予想される。
- コンピュータというものは、Google などのデータセンターがサーバとして、Laptop や Handheld が端末として利用されるものとなってきた。
- 100 以上の core に対して、生産的かつ経済的、ポータブルなソフトウェアの開発を行うために、新しいセンターを設立した。
- アプリケーションとしては、多数の CPU を利用し高性能な計算が要求される音楽、イメージ検索、冠状動脈障害、パレルブラウザなどを扱うこととした。必要な計算要素に分解し、最適化を検討した。
- パレルソフトウェアの開発では、効率の追求と生産性の追求を分離した。効率を追求するのはプログラマーの 10%であり、残りの 90%は生産性を追及しているからである。
- FPGA 上に 12 個の 32-bit RICS Core (90MHz)を構築し、84 個の FPGA を接続した 1008 core の”RAMP Blue”という試験用計算機を構築した。
- 今後、電源やパフォーマンスのボトルネック解析を行い、最適化を図っていく考えである。



Context: Re-inventing Client/Server

- "The Datacenter is the Computer"
- Building sized computers: Google, MS, ...
- "The Laptop/Handheld is the Computer"
- '07: Number HP laptops > desktops
- 1B+ Cell phones/yr, increasing in function
- e.g., Apple iPhone
- Laptop/Handheld as future client, Datacenter as future server

Why might we succeed this time?

- No Killer Microprocessor
 - No one is building a faster serial microprocessor
 - For programs to go faster, SW must use parallel HW
- Multicore Synergy with SaaS / Util. Computing
- All the Wood Behind One Arrow
 - Whole industry committed, so more working on it
- Vitality of Open Source Software
 - OSS community is a meritocracy, so it's more likely to embrace technical advances

Why might we succeed this time?

- Single-Chip Multiprocessors Enable Innovation
 - Enables inventions that were impractical or uneconomical when multiprocessors were 100s of chips
- FPGA prototypes shorten HW/SW cycle
 - Fast enough to run whole SW stack, can change every day vs. every 4 to 5 years when do chips
- Necessity Bolsters Courage
 - Since we must find a solution, industry is more likely to take risks in trying potential solutions

Need a Fresh Approach to Parallelism

- Berkeley researchers from many backgrounds meeting since Feb. 2005 to discuss parallelism
 - Circuit design, computer architecture, massively parallel computing, computer-aided design, embedded hardware and software, programming languages, compilers, scientific programming, and numerical analysis
- Tried to learn from successes in high performance computing (LBNL) and parallel embedded (BWRC)
- Led to "Berkeley View" Tech. Report 12/2006 and new Parallel Computing Laboratory ("Par Lab")
- Goal: Productive, Efficient, Correct, Portable SW for 100+ cores & scale as double cores every 2 years (!)

5 Themes of Par Lab

1. Applications
 - Compelling apps drive top-down research agenda
2. Identify Common Design Patterns and "Bricks"
3. Developing Parallel Software with Productivity, Efficiency, and Correctness
4. OS and Architecture
5. Diagnosing Power/Performance Bottlenecks

Theme 1: Applications. What are the problems?

- "Who needs 100 cores to run N/S Word?"
 - Need compelling apps that use 100s of cores
- How did we pick applications?
 - Enthusiastic expert application partner, leader in field, promise to help design, use, evaluate our technology
- 1. Compelling in terms of likely market or social impact, with short term feasibility and longer term potential
- 2. Requires significant speed-up or a smaller platform
- 3. As a whole, applications cover the most important
 - Platforms (handheld, laptop) Markets (consumer, business, ...)

Compelling Laptop/Handheld Apps

(David Wessel)

- Musicians: Insatiable need for computation
 - More channels, instruments, more processing, more interaction
 - Latency low (5 ms), must be reliable (clicks)
- Music Enhancer
 - Enhanced sound delivery systems for home sound systems using large microphone, speaker arrays
 - Handheld recreates 3D sound over ear buds
- Hearing Augmenter
 - Laptop/Handheld as accelerator for hearing aid
- Novel Instrument User Interface
 - New composition and performance systems beyond keyboards, input device for Laptop

Berkeley Center for New Music and Audio Technology (CNMAT) created a compact loudspeaker array: 10-inch diameter isosahedron incorporating 120 tweeters.

Content-Based Image Retrieval

(Kurt Keutzer)

- Built around Key Characteristics of personal databases
 - Very large number of pictures (>5K)
 - Non-labeled images
 - Many pictures of few people
 - Complex pictures including people, events, places, ...

Coronary Artery Disease (Tony Keaveny)

- Modeling to help patient compliance?
 - 450k deaths/year, 16M symptomatic, 72M High BP
- Massively parallel, Real-time variations
 - CFD FEsolid (non-linear), fluid (Newtonian), pulsatile

Parallel Browser (Ras Bodik)

- Web 2.0: Browser plays role of traditional OS
 - Resource sharing and allocation, Protection
- Goal: Desktop quality browsing on handhelds
 - Enabled by 4G networks, better output devices
- Bottlenecks to parallelize
 - Parsing, Rendering, Scripting
- "SkipJax"
 - Parallel replacement for JavaScript/AJAX
 - Based on Brown's FlapJax

Theme 2: What to compute?

- Look for common computations across many areas
- 1. Embedded Computing (42 EEMBC benchmarks)
- 2. Desktop/Server Computing (28 SPEC2006)
- 3. Data Base / Text Mining Software
- 4. Games/Graphics/Vision
- 5. Machine Learning
- 6. High Performance Computing (Original "7 Dwarfs")
- Result: 13 "Motifs" (Use "motif" instead when go from 7 to 13)

"Motif" Popularity (Red Hot → Blue Cool)

How do compelling apps relate to 13 motifs?

Theme 3: Developing Parallel SW

- 2 types of programmers → 2 layers
- Efficiency Layer (10% of today's programmers)
 - Expert programmers build Frameworks & Libraries, ...
 - "Bare metal" efficiency possible at Efficiency Layer
- Productivity Layer (90% of today's programmers)
 - Domain experts / Naive programmers productively build parallel apps using frameworks & libraries
 - Frameworks & libraries composed to form app frameworks
- Effective composition techniques allows the efficiency programmers to be highly leveraged; major challenge

Machine Learning to Aid Autotuning

- ML to draw inferences from automatically constructed models to automate autotuning
 - Does not rely on application or architecture knowledge
- Learn relationship between set of optimization parameters and resultant performance metrics
 - KCCA finds such multivariate correlations
 - Use relationships to optimize performance-energy
- Stencil: Within 1% expert guided autotune, 4000X faster than exhaustive search

Theme 4: OS and Architecture

(Krzysztof Asanovic, Eric Brewer, John Kubiatowicz)

- HW Solutions: Small is Beautiful
 - Expect many modestly pipelined (5- to 9-stage) CPUs, FPGAs, vector, SIMD Proc. Elmts
 - Offer HW partitions with 1-ns Barriers
- Deconstructing Operating Systems
 - Resurgence of interest in virtual machines
 - Leverage HW partitioning for thin hypervisors
 - Allow SW full access to HW in partition

1008 Core "RAMP Blue"

(Wawrzyniec Krasnow, ... at Berkeley)

- 1008 = 12 32-bit RISC cores / FPGA, 4 FPGAs/board, 21 boards
- Simple MicroBlaze soft cores @ 90 MHz
- NASA Advanced Supercomputing (NAS) Parallel Benchmarks (all class S)
- RAMPants creating HW & SW for many-core community using next gen FPGAs
 - Chuck Thacker & Microsoft designing next boards
 - 3rd party manufacturing and selling boards
 - Gateway, Software BSD open source

Par Lab Domain Expert Deal

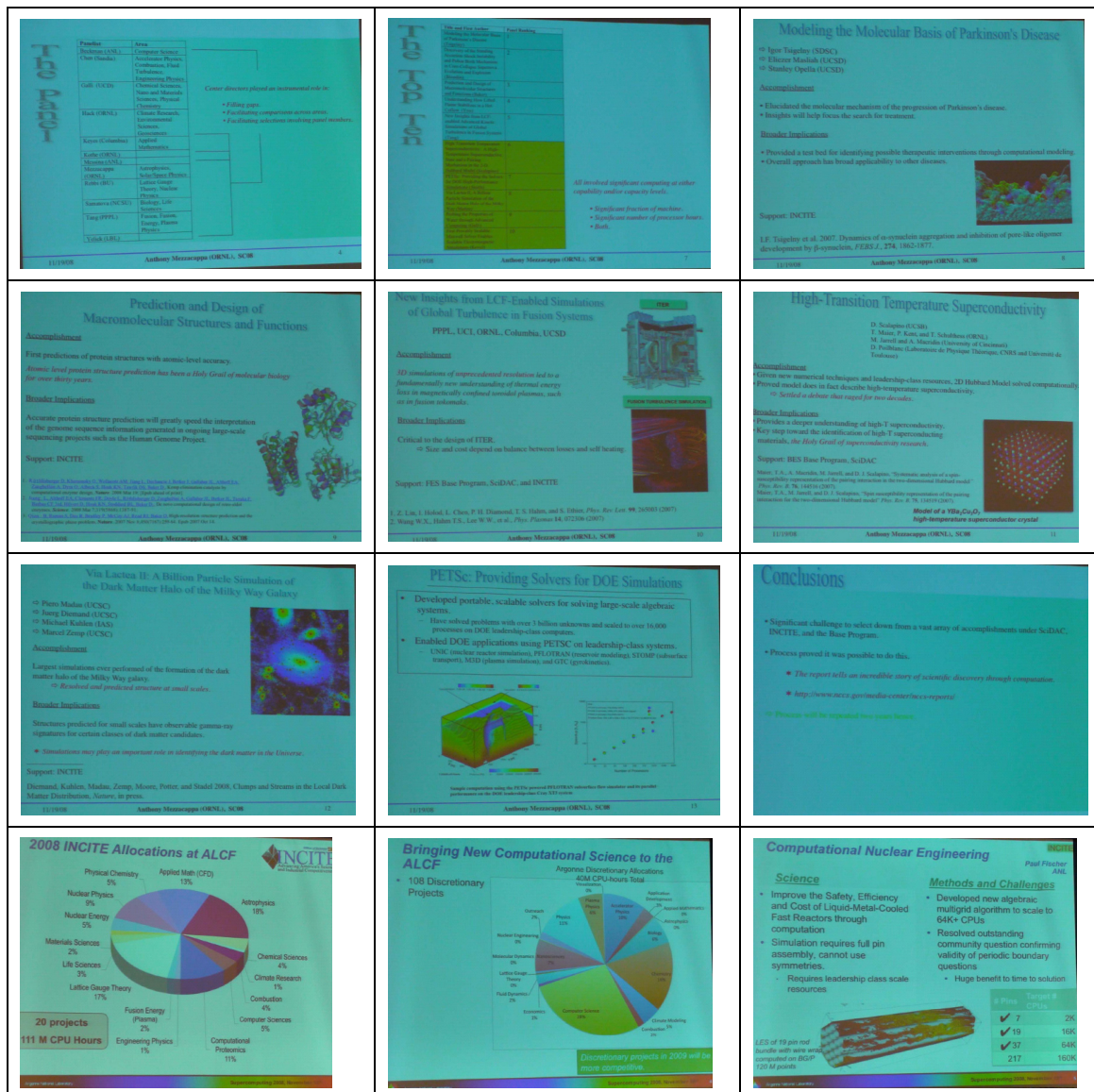
- Get help developing application on latest commercial multicores / GPUs and legacy OS
 - + Develop using many fast, recent, stable computers
 - + Develop on preproduction version of new computers
 - Conventional architectures and OS, but many types
- Will help port app to innovative Par Lab Arch and OS implemented in "RAMP Gold"
 - + Innovate for apps of future (vs. benchmarks of past)
 - Runs 20X slower than commercial hardware

Theme 5: Diagnosing Power/Performance Bottlenecks (Demmel)

- Collect data on Power/Performance bottlenecks
 - Aid autotuner, scheduler, OS in adapting system
- Turn into info to help efficiency-level programmer?
 - Why not using 100% of memory bandwidth?
- Turn into info to help productivity programmer?
 - If I change it like this, impact on Power/Performance?
- A Counter Standard for all multicores?
 - Measuring utilization accurately >>> New Optimization

⑨ BOF 「Petascale Computing Experiences on Blue Gene/P」

実質的には、DOE の INCITE(Innovative and Novel Computational Impact on Theory and Experiment)プログラムなどの下での Blue Gene を用いた計算事例の紹介のみであった。会場から具体的なチューニングの状況について興味がある旨の声が出ていたが、技術的な部分には対応していない発表であった。オークリッジ研究所においては、INCITE プログラムの下で、たんぱく質の MD 計算、ITER の計算、ダークマターに関する粒子系の計算、大規模代数計算などが行われたとのことであった。アルゴンヌ研究所では、CFD や原子力技術をはじめ、20 のプロジェクトに計算資源をアサインしているとのことであった。



⑩ BOF 「User Experiences with Large SGI Altix ICE」

SGI ICE は Xeon ベースのクラスターシステムである。TOP500 の 25 位以内としては、NASA(3 位)、ニューメキシコ計算センター(12 位)、フランス GENCI(14 位)、Total 石油調査(17 位)が導入しているスパコンである。

1 Core あたりのメモリは、2～4GB がトレンドである。2GB で導入したところは、少し苦労しているような感じであった。量子化学計算のように、メモリの大きさがものを言う場合には向いていないという感じであった。安物は安物で苦労しなくてはいけないという感覚は共通のようである。

フランス GENCI のシステムは、32 ラック中 24 ラックを計算用に利用している。設置面積は 63m² で、消費電力は 640kW である。1/4ITER トーラスのシミュレーションでは、1360 億点を 4096Core 上で計算するとのことである。次期 IPCC に利用する計算では、864Core を使い、28km メッシュの 1442×1021×46 点について 300 年間分の計算を行うとのことである。現在、グランドチャレンジとして、燃焼シミュレーションを 8192Core で進めており、12000Core まで引き上げるとのことである。

SGI ICE の利点は、gpgpu や clearspeed を導入するハイブリッド構成が可能なことや、効率、安定性がよいことである。また、ユーザーの満足度は高い。なお、IPoIB のマイナーなバグがあったりするが、改良されるであろう。

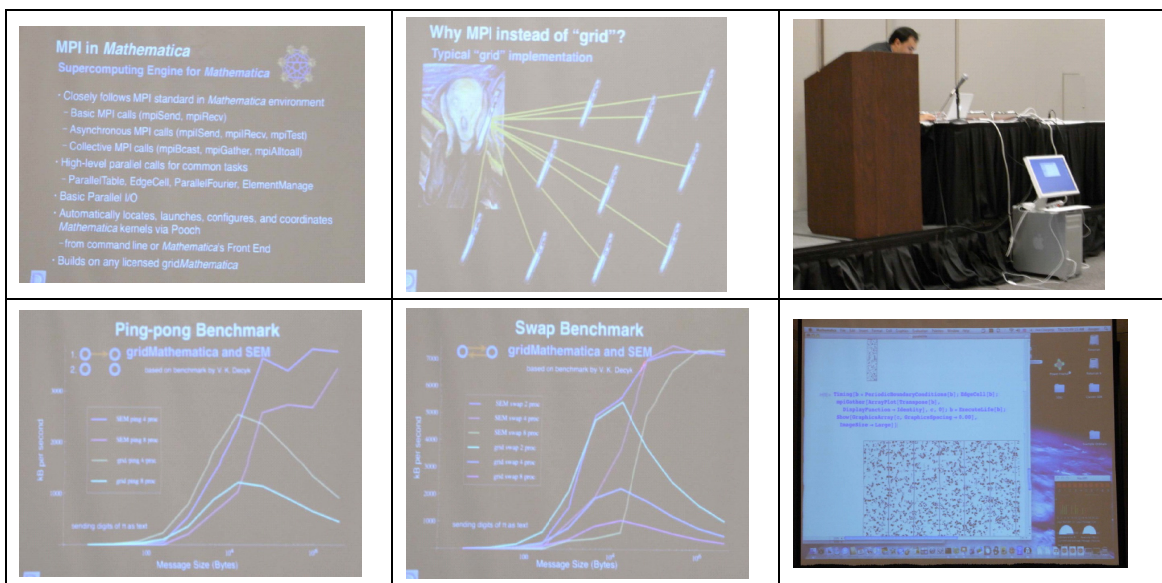
実用の計算では、チープシステムにおいてはファイルの I/O 周りが厳しくなりがちで、そのあたりは仕方ないと考えている様子であった。

⑪ Exhibiter Events 「Supercomputing Engine for Mathematica」 by D. E. Dauger (Dauger Research Inc.)

Matheamtica 上で MPI を利用できるための環境構築に関する発表であった。MPI プログラミングを Mathematica でそのまま行うことができるライブラリとみなすことができる。Mathematica 上で大規模計算を通常しないとは考えられるが、初学者にとって役立つこともあることから今後の発展が期待される。

発表のポイントは次の通りである。

- ・ 既存製品の **gridMathematica** では、データが大きくなると使い物にならなくなるとのことである。**grid** タイプの計算では、データを 1 つの **core** に集約することとなるため、マスターノードで処理できなくなるためである。
- ・ その対応として、MPI による **Domain Decomposition** の計算を実現すべく、MPI の **Send**、**Recv**、**Gather** などの機能を備えた環境を整備した。
- ・ **gridMathematica** のライセンス上で動作するように構築されている。
- ・ 具体的には、**Pooch** という **Dauger Research** の MPI に関連するソフトを用いて複数の **Mac** でクラスターを構成した環境において、**Mathematica** の **FrontEnd** から MPI と同じような機能をするコマンドを発行することで MPI 並列計算を行うことができる。
- ・ セルオートマトンやプラズマダイナミクスのシミュレーションについて、領域分割した計算が可能である。
- ・ **gridMathematica** と比較し、実際に大規模な計算が可能である。



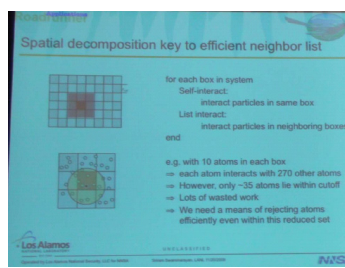
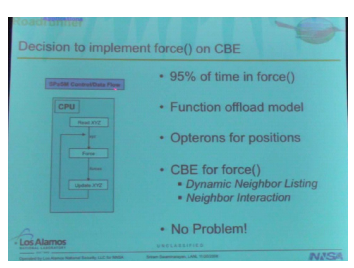
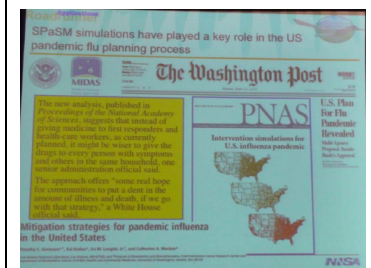
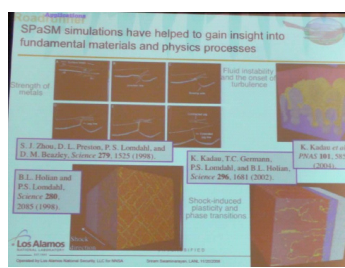
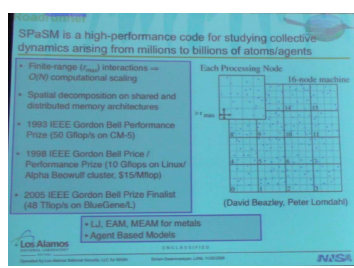
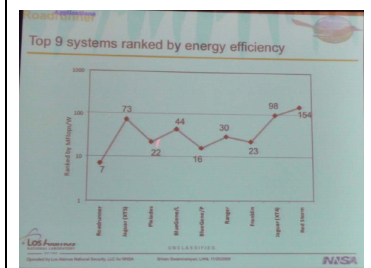
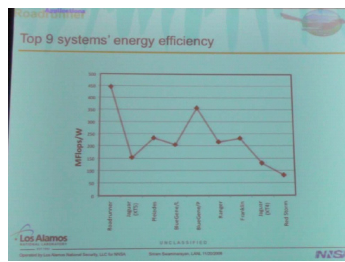
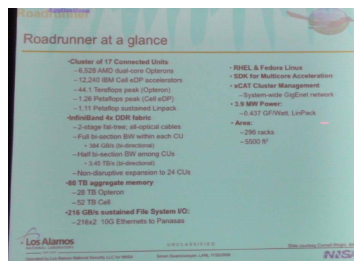
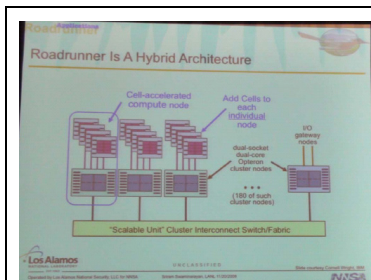
⑫ ACM Gordon Bell Finalist 「369 Tflop/s Molecular Dynamics Simulations on the Roadrunner General-purpose Heterogeneous Supercomputer」 by Sriram Swaminarayan et. al. (Los Alamos)

Roadrunner 上で MD コードを最適化した研究の発表であった。Cell の MD プログラムのチューニングをしっかりとやっている印象を受けた。また、個人的には、Cell-SPE と Opteron の間の通信が多いような気がしたが、実際にやってみないと判断がつかないと思われる。

本発表の論文の PDF は、<http://scyourway.nacse.org/conference/view/gb124> を通じ IEEE のアーカイブや ACM のデジタル図書館から取得できる。

発表のポイントは次の通りである。

- Roadrunner は、17 のユニットに、6528 個の Dual-core Opteron と 12240 個の Cell を搭載した計算機である。Opteron のピークは 44.1 テラ flops で、Cell が 1.26 ペタ flops である。ユニット内では 384GB/s の InfiniBand で接続され、ユニット間は 3.45TB/s の InfiniBand で、24 ユニットまで接続できる。ファイルシステムとは、216GB/s の I/O が可能である。
- Roadrunner は Mflops/W の値が圧倒的に低い、エコなシステムである。
- SPaSM(Scalable Parallel Short-range Molecular dynamics)は、1990 年代にロスアラモスで開発された領域分割型の並列 MD コードである。2005 年に Blue Gene/L 上で動作するように移植されている。SPaSM は、金属の計算で利用されるとともに、感染症のエージェント計算にも用いられている。
- 原子間の力の計算では、領域分割を用いるのが、(当然) 最も効率的である。
- スタンドアロンの Cell での MD は、Opteron の 3.5 倍の速度であった。
- Cell は、ベクトル向けのデータ構造である必要があり、データ構造の見直しを行った。SPaSM の通信は、1timestep でプロセスあたり 50000 メッセージとなることがある。今回は、6 までに通信を減らした。このようにしても、通信とデータの入れ替えに 20%の時間を使っている。
- Opteron は通信を担当し、Cell の SPU で他の処理を行った。
- 結果としては、Opteron の 6~9 倍まで高速化することができた。
- 平衡系の問題やエージェントシミュレーションなどでは、単精度にすることが可能である。その場合さらに、多くの flops 値を稼ぐことができる。8 ユニットの場合、倍精度では 195 テラ flops であるが、単精度では 471 テラ flops となる。



Optimizing the neighbor listing

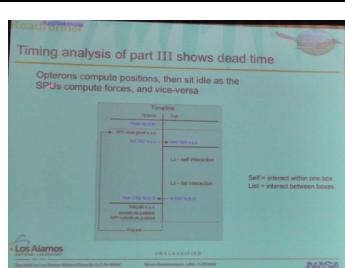
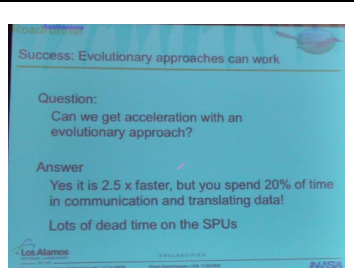
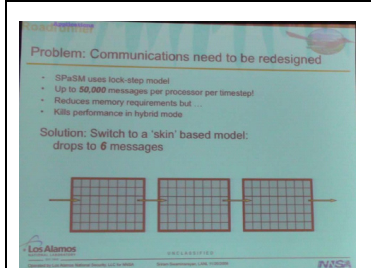
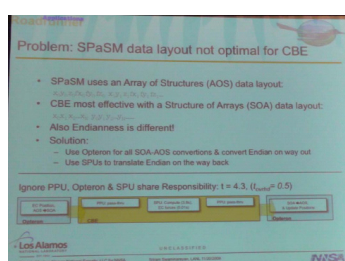
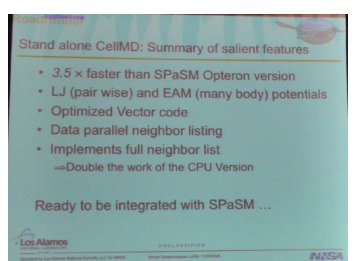
Method	Naive	Simple	Data Parallel
% time in NL	60%	60%	25%

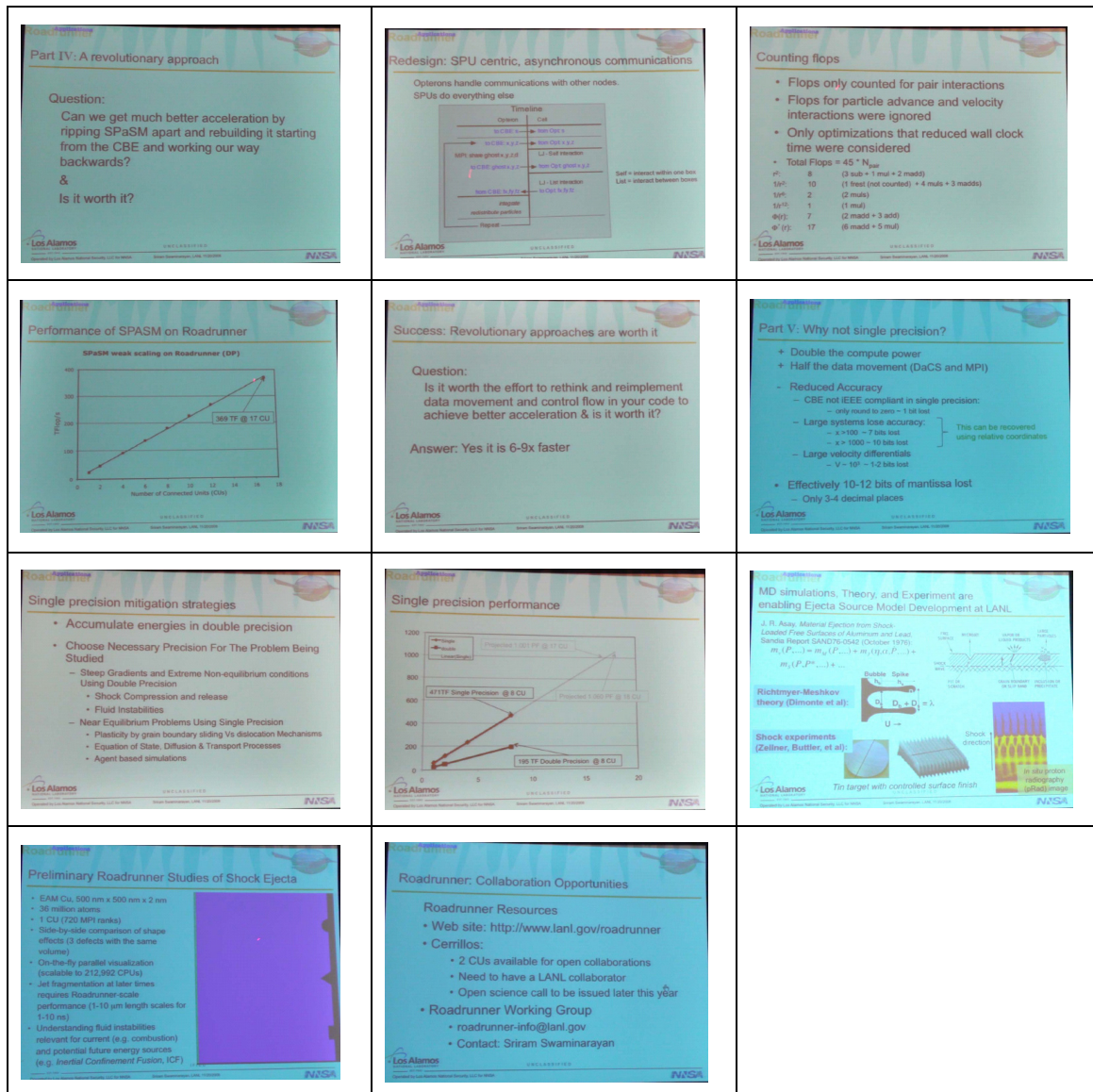
Los Alamos

UNCLASSIFIED

Source: Los Alamos National Laboratory, LANL-11-00000

ROADRUNNER





⑬ BOF 「The HPC Advisory Council Initiative」

この BOF は、実質的には、60 を超えるベンダー企業、大学との HPC に関するコンソーシアムのミーティングであった。HPC 技術、人的ネットワーク、教育、ベストプラクティスなどを共有している。この BOF では、下記のことについて、フリーディスカッションがなされた。特筆すべき意見は見られなかった。直感的には、お付き合い程度で、機能していないコンソーシアムのように感じられた。

- HPC 利用の上での最大の弱点は？
- ベンダーとエンドユーザーの関係のレベルは？
- HPC リソースをちゃんと使っているのか？
- 貴社は、本コンソーシアムの人的にネットワークにどういう経験について情報提供できるか？
- よりよい生産性のために HPC を使いたがるユーザーに思いとどませた事例はあるか？
- 生産性か効率か？違いはあるのか？
- 本コンソーシアムでの活動経験は？

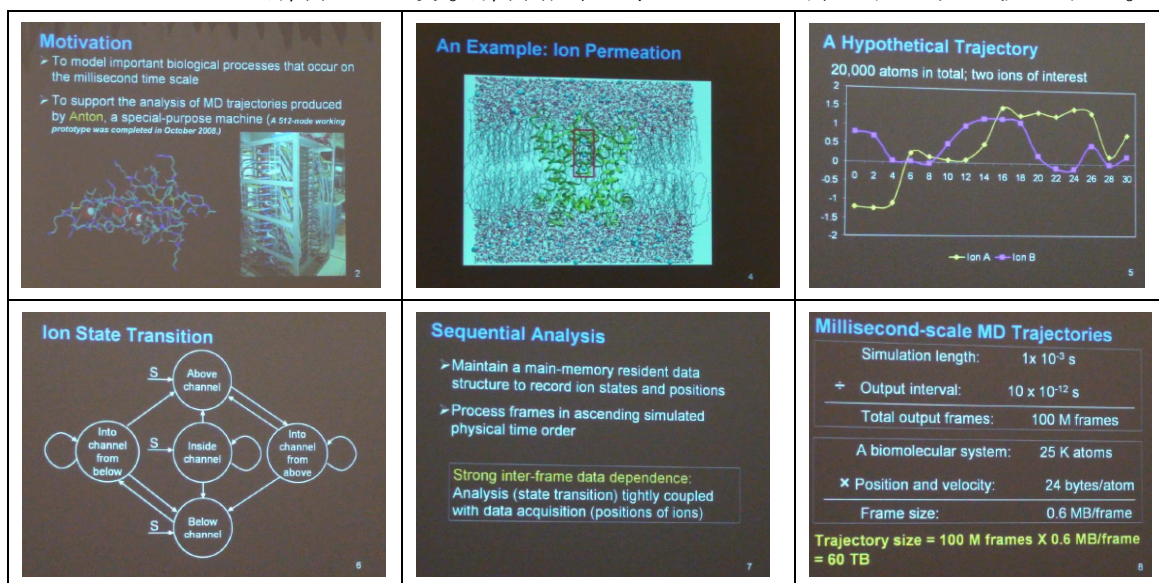
⑭ Papers 「A Scalable Parallel Framework for Analyzing Terascale Molecular Dynamics Trajectories」 by Tiankai Tu et. al. (D. E. Shaw)

MD 専用計算機 Anton の稼動後に予想される膨大なデータについて、データ解析を並列に行う手法の検討である。2 万原子の中から 2 つのイオンの位置座標の時系列データを取り出す手法を提案している。提案の手法は分散入力ポイントであり、性能はファイルシステムに依存するところが大きいと考えられる。出力はファイル出力ではなく、可視化アプリのメモリバッファへの出力を実装している。手早く計算結果を見るためには、有用なツールである。

論文の PDF は、<http://scyourway.nacse.org/conference/view/pap318>。

発表のポイントは次の通りである。

- Anton は、512 ノードのプロトタイプが 2008 年 10 月に稼動している。
- 膜タンパクでのイオンの動き（軌道）に注目していることから、約 2 万の原子の時系列データから特定のイオンの時系列データを取り出す必要がある。
- 10 ピコ秒ごとにスナップショットの出力を想定し、1 ミリ秒の長さでは、1 億毎のスナップショットで 60TB のデータサイズとなる。ローカルディスクのバンド幅を 100MB/s とすると、読み込みに 1 週間以上かかる。
- HiMACH は、Google type の Map Reduce による処理を念頭においた、MD 軌道を解析する Python プログラムである。フレームごとのデータの取得を行い、その後、通信により、フレーム間のデータ解析を行う手法である。
- 512Core を用いると、Single Core に比べ 100 倍以上速くなる。
- 2 万原子の中から 2 つのイオンの位置座標の時系列データを取り出す場合、1TB のデータ解析が 15 分。解析結果は、VMD の内部バッファへ投入する。



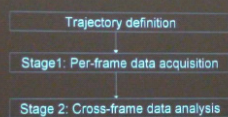
Problems with Sequential Analysis

Millisecond-scale trajectory size : 60 TB
 Local disk read bandwidth: 100 MB / s
 Time to fetch data to memory: > 1 week
 Analysis time: Varied
 Time to perform data analysis : Weeks

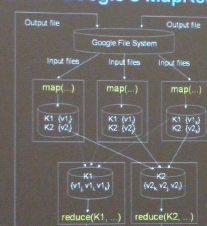
Sequential analysis lack the computational, memory, and I/O capabilities!

A New Data Analysis Model

- Specify which frames to be accessed
- Decouple data acquisition from data analysis



Inspiration: Google's MapReduce



Trajectory Analysis Cast Into MapReduce

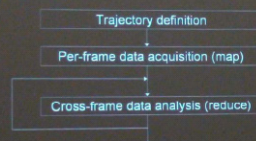
- Per-frame data acquisition (stage 1): map()
- Cross-frame data analysis (stage 2): reduce()
- Key-value pairs: connecting stage1 and stage2
 - Keys: categorical identifiers or names
 - Values: including timestamps
 - Examples: (ion_id_i, (t_k, x_{jk}, y_{jk}, z_{jk}))

HiMach

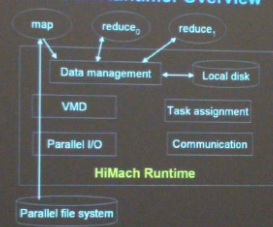
- A MapReduce-style API that allows users to write Python programs to analyze MD trajectory
- A parallel runtime that executes HiMach user programs in parallel on a Linux cluster automatically
- Performance results on a Linux cluster:
 - 2 orders of magnitude faster on 512 cores than on a single core
 - 1 TB data analyzed in 15 minutes

Chained Reduce Operations

- Support multiple rounds of analysis
- Exploit a Python language feature (the yield statement)



The HiMach Runtime: Overview



Task Assignment

- Assumption: equal-weighted tasks
- Partitioning of the workload:
 - A block data decomposition algorithm
 - No inter-processor communication when assigning map tasks
 - One collective operation (MPI_Allgather) when assigning reduce tasks

Parallel I/O

- Assumption:
 - Frames organized as megabyte-sized files
 - Files stored in NFS or in a parallel file system
- How to read the frames:
 - No communication needed to coordinate I/O
 - Reads issued by HiMach
 - Frame data passed to user-defined map()
- How to write output results:
 - Writes issued by user programs directly

Interaction with VMD

- VMD: Visual Molecular Dynamics (UIUC)
- Run an instance of VMD on each core
- Load molecule structure files
- Push frame into VMD's internal buffer
- Support frame caching and prefetching

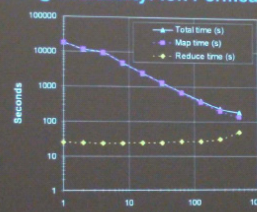
Experiment Setup

- Cluster configuration:
 - A commodity Linux cluster with 128 nodes
 - Two Intel Xeon 2.66 GHz quad-core processors per node
 - Two 250 GB SATA disks organized in RAID 1 per node
 - CentOS 4.6 with a Linux kernel of 2.6.22
 - Nodes connected via a Gigabit Ethernet connection
- Storage configuration:
 - PVFS 2.7.0 (a parallel file system developed at Argonne National Laboratory) on 64 nodes
 - Both the I/O server and the metadata servers on all the 64 nodes

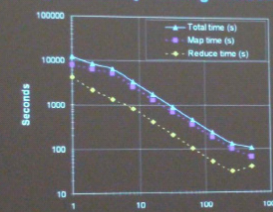
Experiment Setup

- Driving applications:
 - Ion permeation: Time series analysis (154 GB)
 - Sliding-window RMSD: Measurement of structural "proximity" between two frames (85 GB)

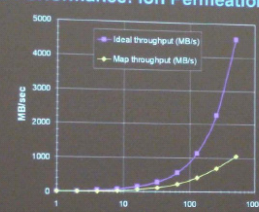
Strong Scalability: Ion Permeation



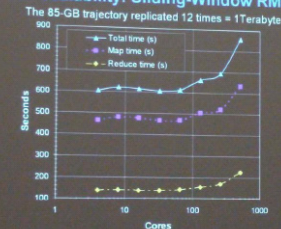
Strong Scalability: Sliding-Window RMSD



I/O Performance: Ion Permeation



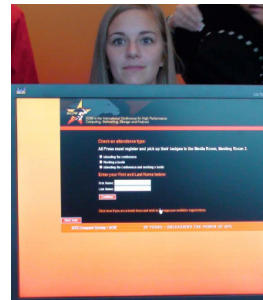
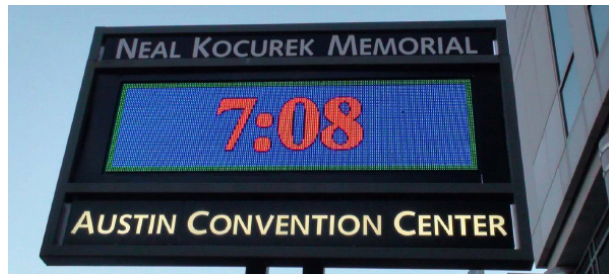
Weak Scalability: Sliding-Window RMSD



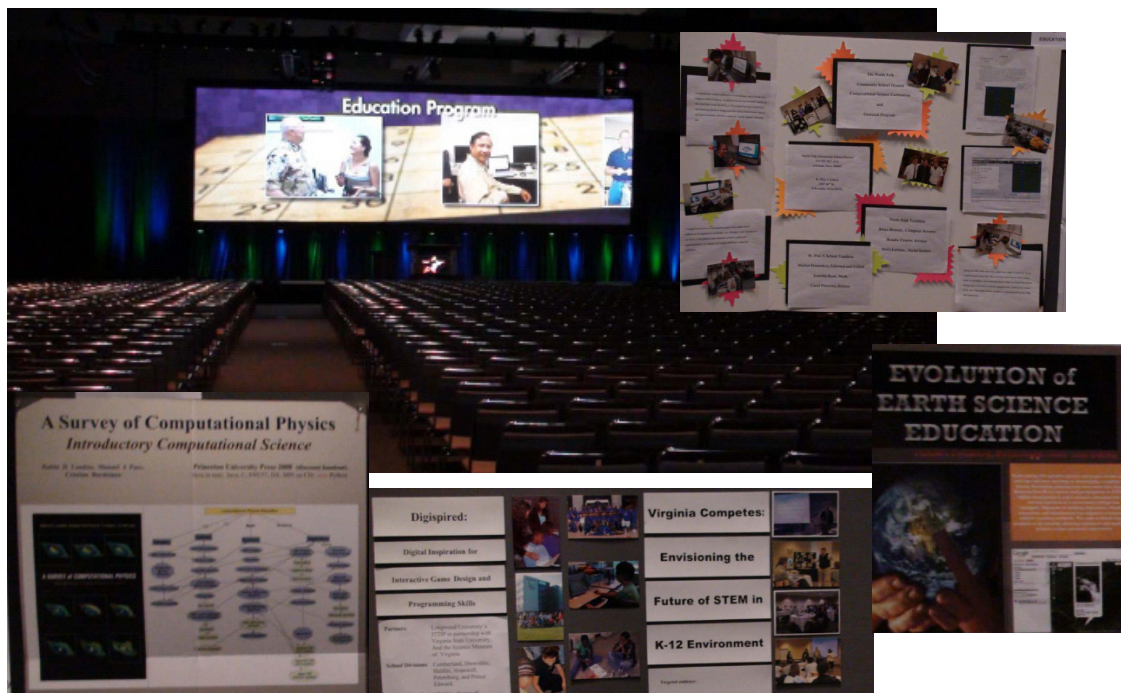
Summary

- HiMach = MapReduce + MPI + Python + VMD
- A new MD trajectory parallel analysis model
- A new interface for defining trajectories
- A novel method of running VMD in parallel
- An extension to MapReduce to support chained reduce operations
- Potentially applicable to data analysis for other types of simulations

● SC08 の会場へ



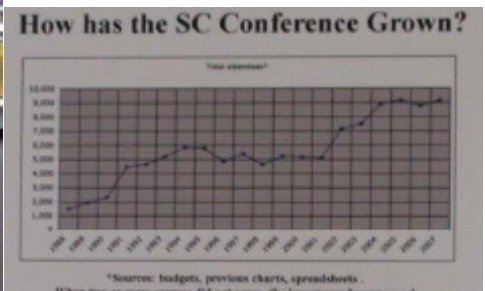
● Education Program

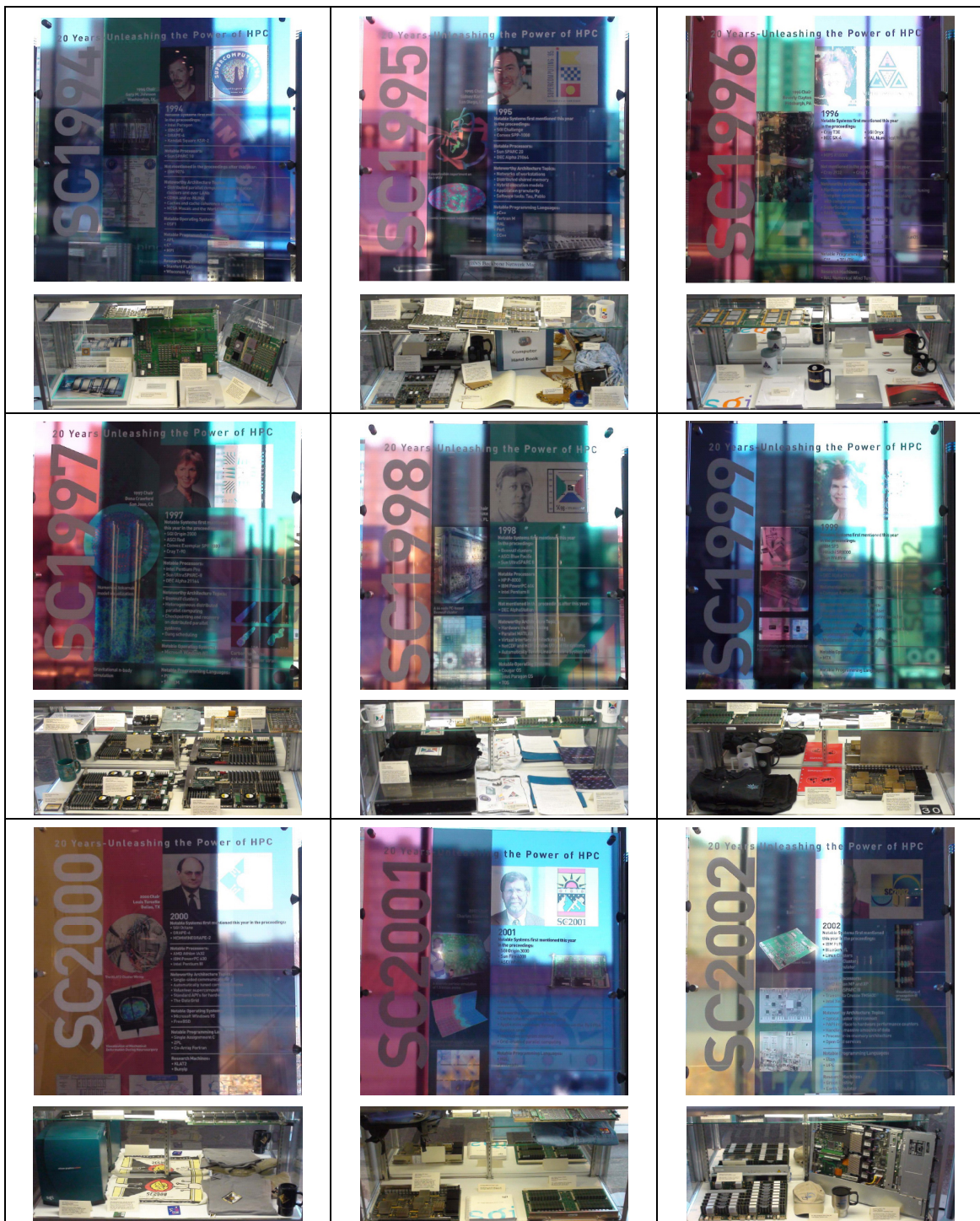


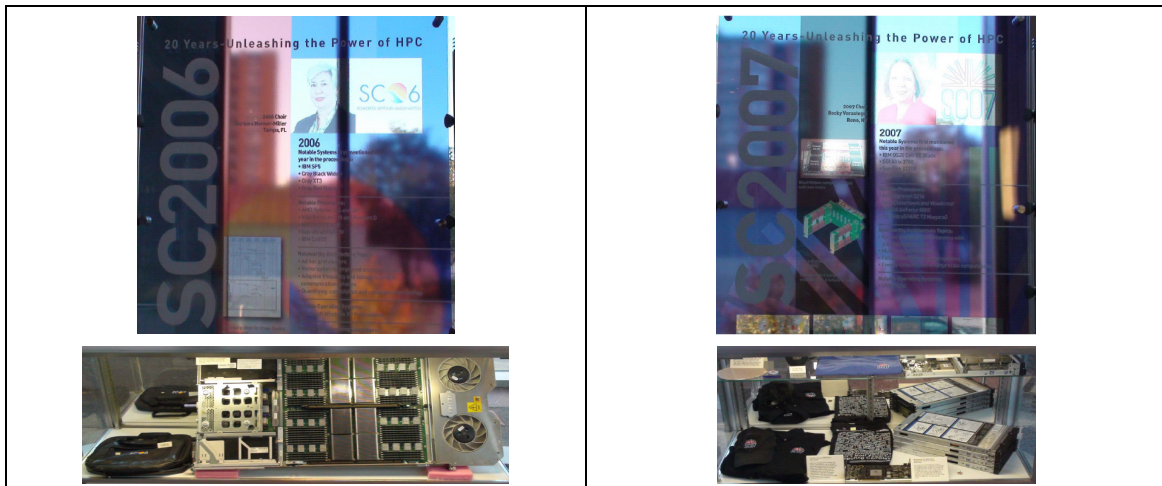
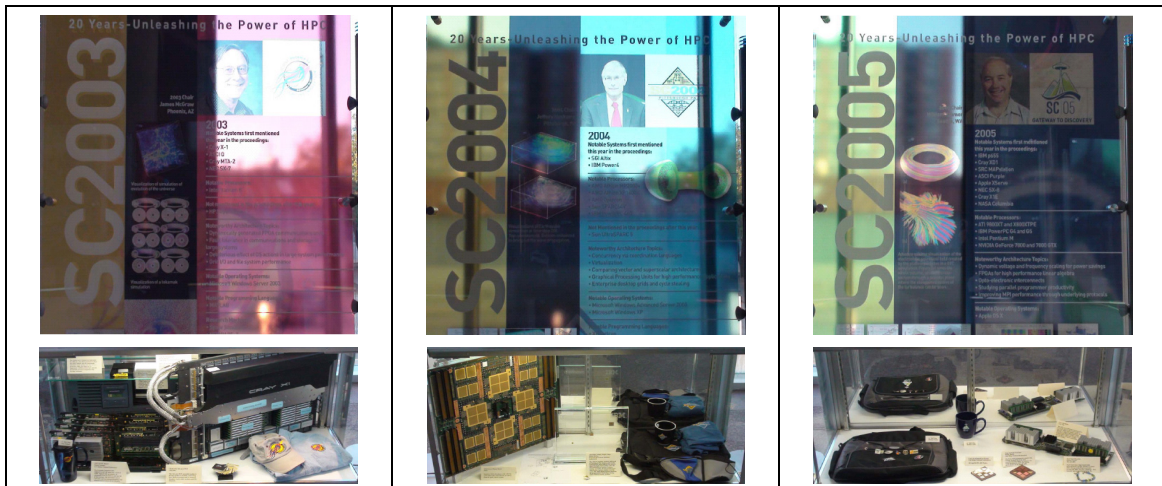
● SCを支えるネットワーク



● SC20 年の歴史

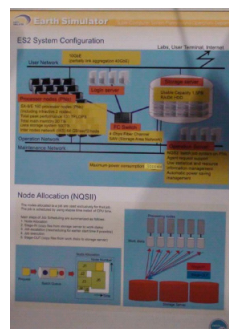
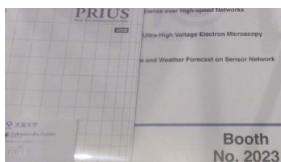
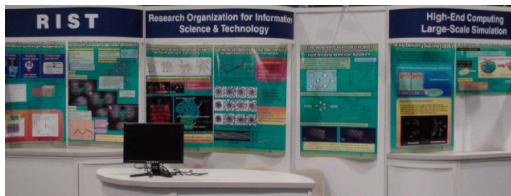


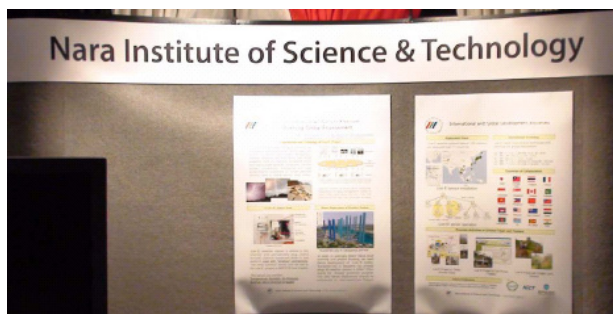
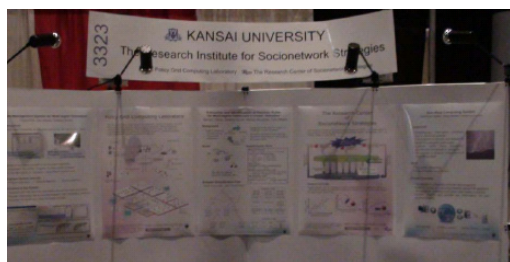




● 展示会場 （日本コレクション）







- ペタ flops 級の計算機を持つ大御所のブース





● 代表的なベンダーのブース



● おまけ

