

19th I S C A (Annual International Symposium on Computer Architecture)

参加報告

佐藤三久 (電子技術総合研究所)

1992年5月19日から21日の3日間にわたってオーストラリアゴールドコーストにおいて開催された第19回コンピュータアーキテクチャ国際シンポジウムの模様について報告する。

Report on 19th Annual International Symposium on Computer Architecture

SATO MITSUHISA
(Electrotechnical Laboratory)

I report on 19th Annual International Symposium on Computer Architecture held at Gold Coast, Australia on May 18-22, 1992.

19th ISCA (Annual International Symposium on Computer Architecture) 参加報告

佐藤三久 (電子技術総合研究所)

5月19日から21日の三日間にわたって、オーストラリア・ゴールドコーストで開催された第19回コンピュータアーキテクチャ国際シンポジウム (ISCA) に参加したので、その報告する。シンポジウムはゴールドコーストのラマダホテルにおいて開催された。参加人数は200人前後 (当日にはこれ以上いたようであった) と盛況であった。

プログラム委員会の報告によると、173の投稿論文の内、39の論文が採択され (採択率23%)、2つのパラレルセッションで論文発表が行われた。論文の採択は、北米、ヨーロッパ、日本からバランスを考えて採録されていたようであった。

日本からは、平田ら (松下メディアラボ) の "An Elementary Processor Architecture with Simultaneous Instruction Issuing from Multiple Threads", 筆者ら (電総研) の "Thread-based Programming for the EM-4 Hybrid Dataflow Machine", 清水ら (富士通) の "Low-latency Message Communication Support for the AP1000" の3件であった。

詳しくは論文集を参照して頂くことにして、シンポジウムの模様について報告することにする。報告を行うに当たって、筆者個人のバイアスがあることをあらかじめお断りしておく。

招待講演

keynote address は、ヒューレットパッカートの Bob Raw が HP-PA について、招待講演は、RS/6000 のワンチップ版 (講演者は失念した)、John Hennessy による "The MIPS R4000: Architecture and Performance", Jim Smith (Cray Research) による "What Makes Supercomputers 'Super'", 最後に DEC の Alpha についてであった。Cray の講演を除いては、ワークステーションのいわゆる Killer Micro についてである。

現在のところ、最先端のラインアップとしてはテクノロジーは 0.7 ~ 0.8 ミクロン、内部クロック 100 MHz (外部 50 MHz)、性能は 100 SPECmark 前後というところである。(HP PA-RISC は、既に 120 SPECmark のチップを発表しており、single chip RS/6000、R4000 も同程度性能を達成している。) アーキテクチャ面からいえば、パイプラインの段数が多少変化していることと R4000 のように 2 level Cache を使っていること、さらに 64 ビットのアドレス空間をサポートしていることが大きな点として上げることができるが、むしろ、クロックのスピードを上げることにに関して、テクノロジーとのリンクが最大の問題点であると言えるだろう。

既に話題になっているが、Alpha は DEC の 90 年代のアーキテクチャとして、現在最も早い 200MHz のクロックを達成している。マシンのシリーズの名前が VAX になることからわかるように、これまでの VAX のコードをアセンブレベルでトランスレートすることができ、VMS など移植可能なように設計したということであった。これらは、PAL (Privileged Architecture Library) call と呼ばれるマクロ命令のような機構でサポートされていて、DEC らしい戦略といえるのではないか。テクノロジーの点からチップ写真で目を引いたのは、クロック分配のための大きなトランジスタが背骨のようにおかれていたことである。これからはクロックのスキューを抑え

る技術が大きな問題になるであろう。

Hennessy は MIPS の最新チップである R4000 について講演した。Hennessy によれば、プログラムサイズ（もちろんデータも含むと思うが）は年 1 ビットから $1/2$ ビットのペースで増えており、これからは 64 ビットのアドレス空間が必要であることを強調していた。チップは、UNIX オペレーティングシステムのレベルのプログラムが動くことをシミュレーションで確認してからつくるということであった。Hennessy はスタンフォード大の教授でもあるが、MIPS 社の企業人でもある。Chip のデザインレベルでは R4000 には常時 30 人程度の人に関わっており、結論として高速なチップに必要なものは、いろいろな技術の妥協点を探ることと開発に対するコストをあげていた。製品のラインアップとして、速度は同等で消費電力を下げたバージョンなども準備されているようである（TS と呼ばれていた）。示された OHP には、High end の製品として、1994、1995 年当たりで 0.3 ミクロンで内部クロック 250MHz のものが計画として示されていた。他のチップメーカーも同じような開発ベースで RISC チップを出してくるようである。

J. Smith の supercomputer の講演についてはとくに目新しいものはなく、ベクター処理による高速数値計算とそれをサポートする贅沢なメモリ構成と贅沢な I/O 能力が Super ならしめているということである。

ちなみに、今年の Eckert-Mauchly Award は MIMD や SIMD などの用語で有名な Michael Fynn であった。

パネル

さて、招待講演は Hot Chips なみの Killer Micro のオンパレードであったが、パネルは、アプリケーションよりの人たちによる "What Problems Can Truly Justify Building a Million Processor Machine" とアーキテクチャよりの "What Should the Architecture Be for the Processor Used in a General Purpose Teraflops Computing System" という超並列と Teraflops のおなじみの話題 2 つであった。

詳しく解説するだけの能力はないが、超並列のアプリケーションといってもこれまでのスーパーコンピュータの応用から大きく離れるようなものはないという印象であった。"Justify" という言葉から連想されるように、会場からの反応は市場やそのようなシステムの価格のようなことに終始していたようである。

後者のパネルでは、MIT の Arvind や Hennessy らがパネラとして登場し、MIT のマルチスレッドマシンの *T やスタンフォードの DASH が取り上げられていた。ここでも、誰がどのように使うのかということになるが、こちらでは超並列の multi-user Operating System (time sharing を含む) や I/O がどうなるかについて、会場からの質問があった。筆者の興味を引いた議論として、Teraflops を目指してプロセッサ数が増えていくと、キャッシュなどの局所性を活かす技術が効かなくなっていくがその点はどうするのかという質問があったが、Hennessy は共有メモリのプロセッサでもコンパイラやいろいろな技術を使って局所性は活用できるという強引な議論を展開していた。彼らのアプローチはあくまでも R4000 などの高性能のチップを使って共有メモリのマシンで進んでいこうというようなことなのだが、数人の人たち（特にデータフロー関係の人たちであったが）に後で聞いてみたところ、100 台レベルにしか通用しないのではないかというのが、大方の意見であった。ちなみにパネルで teraflops にするには 100 GFLOPS のプロセッサが 10 台あれば、いいのだという冗談があったが、そうならばたしかに共有メモリで十分であろう。

論文発表

今回のシンポジウムで目についたのは、Multi-thread と名前がつくセッションが2つあり、超並列マルチプロセッサに向けてマルチスレッドアーキテクチャについての発表が多かったことである。

マルチスレッドアーキテクチャでは、オペレーションのレーテンシ、特にリモートのメモリに対するアクセスのレーテンシを隠すアーキテクチャとそのコンパイラ技術、そして実際の性能がどうなるかが焦点になっている。マルチスレッドアーキテクチャのセッションでは、日本からの2件の発表とMITの*Tの発表があった。*TはMITとモトローラの共同開発でモトローラのRISCチップを基にしたマシンであり、超並列アーキテクチャを目指している。また、MITのDallyのグループからは、実行時のスケジューリングを考慮して、コンパイラでスレッドを分割するアーキテクチャについて発表していた(“Processor Coupling: Integrating Compile Time and Runtime Scheduling for Parallelism”, S. W. Keckler and W. J. Dally)。“Improved Multithreading Techniques for Hiding Communication Latency in Multiprocessor”, B. Boothe and A. Ranadeでは、リモートメモリアccessのレーテンシを隠すアーキテクチャの検討が行なわれている。この論文では長いレーテンシを持つ共有メモリを仮定している。

ISCAと併設されて開催されたDataflow Computing Workshopでも見られたことであるが、最近のデータフローマシンの研究は既にマルチスレッドのアーキテクチャにシフトしており、fine-grainのデータフローマシンの研究でもいかに局所性の活用するかが焦点になっている。このような研究の1つとしては、“An Analysis of Loop Latency in Dataflow Execution”, W. A. Najjar, W. M. Miller, A. P. W. Bohmもデータフローのfine-grainのデータフロー計算をいかにclusteringして、実行するかについての発表があった。

これからの汎用の(超)並列マシンに向けて、アーキテクチャとして大いに期待されるマルチスレッドアーキテクチャであるが、同時にコンパイラとアプリケーション、言語が今一つ定まっていなような気がする。ソフトウェアの未熟ゆえ、筆者自身も含めて評価が初歩的な段階に留まっている印象である。Dataflow WorkshopではIdやSISALでのアプリケーションの蓄積が行なわれていて、それを使って多くの評価が行なわれていたが、超並列のアーキテクチャの発展にはもちろんそのアーキテクチャを活かすソフトウェアがさらに重要になっていくのではないだろうか。

一方、共有メモリマルチプロセッサ、あるいは共有メモリのキャッシュの発表は、スタンフォードの発表が多いが目についた。アーキテクチャ的には非常に堅実な路線であり、定量的な比較研究には感心させられる。スタンフォードでは、SPLASHと呼ばれる共有メモリ向けのベンチマークセットが揃っており、評価ソフトウェアもしっかりしている。ちなみに、DASHは16台が稼働しており、64台までのシステムの評価結果が発表されていた。しかし、このままのアプローチが数千台の規模に有効であるかは疑問が残るところである。

メッセージパッシングに関して、ユーザモード実行できるオーバーヘッドの少ないメッセージ通信機構の発表が2件あった。1件は日本からのAP1000の発表であり、もう一件は“Active Messages: a Mechanism for Integrated Communication and Computation”, T. Eicken, D. E. Culler, S. C. Goldstein and K. E. Schauerである。後者ではCM-5などでオーバーヘッドの少ない通信とそれをTAM(Threaded Abstract Machine)と呼ばれる比較的fine-grainの計算モデルと統合して使おうというものである。彼らは、現在あるマシンをマルチスレッドマシンとして使うことを検討しているようである。いずれにしても、台数が多くなりfine-grainになると、メッセージ通信のオーバーヘッドが焦点になるのは当然の成行きであろう。

最後に並列度の限界に関する、同じようなシミュレーションによる研究が2件あったので、それらを紹介する。

“Limits of Control Flow on Parallelism”, M. S. Lam and R. P. Wilson では、control flow による並列性について考察し、複数の control flow が許される場合、さらに Speculative な実行が許された場合の並列度について報告している。従来の 1 つの Control flow しか許されない Super Scalar や VLIW では従来の SPEC benchmark の gcc や latex など逐次プログラムでは並列度が 2 ～ 3 程度であることは良く知られているが、control flow を複数用いることで 5 ～ 10 程度、さらに speculative 実行することで 20 ～ 200 程度になる。データフローマシンなどでは複数の control flow を実現することができるが、論文では、speculative 実行が大幅な並列実行には必要であることを強調している。そのためのオーバーヘッドが少ない有効なアーキテクチャが見つかるかは今後の課題であろう。

また、“Dynamic Dependency Analysis of Ordinary Program”, T. M. Austin and G. S. Sohi では、命令トレースからデータ依存関係だけで計算できる場合の並列度、すなわち、レジスタやメモリが完全に renaming できた時の並列度について論じている。このデータは命令トレースからとったもので、どの様なアーキテクチャをもってしても完全に renaming できるわけではない。しかし、レジスタの renaming で結構並列度が向上することを強調している。(pure data flow machine はある意味では renaming しているのであるが、それとこのシミュレーション結果の関連はあるのであろうか?)

これらには全てが予言できる状態でのシミュレーション (oracle) の結果がしめされており、gcc などの逐次プログラムでは並列度は数百程度である。当然のことながら、数値計算などのプログラムでは両者の報告でも大きな並列度が得られることが示されている。

第 20 回 ISCA

なお、来年の第 20 回 ISCA は、San Diego, USA で開催予定であり、論文投稿の期日等は例年通りである。