

統計データ解析における並列計算機システムの 性能評価

筑波大学* 下 平 文 彦

筑波大学** 白 川 友 紀

統計数理研究所 田 村 義 保

(受付 1998 年 1 月 22 日; 改訂 1998 年 6 月 30 日)

要 旨

近年、情報処理機器の発展に伴い社会のあらゆる場所で大量のデータが蓄積されるようになってきており、その解析が重要になってきている。しかし、大量のデータに対して精密な解析を行なおうとすると、膨大な時間がかかってしまう。

本研究では、並列計算機上で統計データ解析の高速な処理を行なうことを目的とする。パソコン・クラスタ、並列計算機 CP-PACS 及び IBM RS/6000 SP 上で実行し、その処理時間、並列処理効率、及び速度向上の様子などを測定し、各計算機のデータ解析分野への適応性などを評価した。尚、本研究では基本統計量、重回帰分析、主成分分析について測定した。

その結果、100 変量、10 万サンプルのデータを 16 台の PU を使用して処理する場合には、基本統計量、重回帰分析、主成分分析の処理効率は CP-PACS ではそれぞれ、100%以上(463.6%)、30.4%、78.7%となった。同様の条件の時、RS/6000 SP では、98.4%、23.4%、100%以上(137.3%)となった。また、8PU のパソコン・クラスタでは 100 変量、1 万サンプルのデータの時、それぞれ 87.1%、3.4%、54.8%となった。

実処理時間を比較すると基本統計量では CP-PACS が、通信の頻繁に行なわれる重回帰分析では RS/6000 SP が、主成分分析では RS/6000 SP が最も適しているという結果が得られた。また、主成分分析では計算途中に通信が行なわれないこともあり、パソコン・クラスタでも実用に耐え得る速度が得られた。

キーワード：並列計算機，CP-PACS，RS/6000 SP，パソコン・クラスタ，並列処理効率。

1. はじめに

近年、情報処理機器の発展に伴い社会のあらゆる場所で大量のデータが蓄積されるようになってきており、その解析が重要になってきている。

データの集まりを捉え、一つ一つのバラバラなデータからは得られない、データの集団が持つ様々な特徴（法則）を眼に見えるような形で捉えようとするのが最近のデータ解析の特徴である（村上・田村（1988），渡部 他（1985））。また、このような考え方に基づいて、データベースの分野ではデータマイニングという、データ集団の特徴を抽出する研究も盛んに行なわれ始

* 工学研究科：〒305-8573 茨城県つくば市天王台 1-1-1.

** 構造工学系：〒305-8573 茨城県つくば市天王台 1-1-1.

めている (Agrawal et al. (1993)). しかし, 大量のデータに対して, データのあらゆる組合せについての解析のような精密な解析を行なおうとすると, 膨大な時間がかかってしまう. そこで, 近年, より一般化してきた並列計算機のような高度な計算機システムを使用することが期待されている.

以前, 我々は, 並列計算機 QCDPAX 上での統計データ解析の並列処理とその並列処理効果の評価を行なった. しかし, この研究では QCDPAX 専用の言語を用いて開発を行なったために, 汎用性がなかった (下平 他 (1995)).

本研究では, より汎用性のあるプログラムを開発することを念頭におき, 現在広く使われている, メッセージ通信用の標準ライブラリである MPI (Message Passing Interface) を用いてプログラム開発を行なった. MPI を使用して開発することにより, プログラムの移植性 (可搬性) が高くなり, 異なる環境での処理の実行が容易になった.

こうして開発したプログラムを, パソコン・クラスタ, 並列計算機 CP-PACS, 及び IBM RS/6000 SP 上で実行し, その処理時間, 並列処理効率, 及び速度向上の様子などを測定し, 各計算機のデータ解析分野への適応性などを評価した.

2. 並列計算機

並列計算機とは, 複数台のプロセッサを相互結合網を介して並列に結合して, 全体として高い処理能力を発揮させることを目的とした計算機システムである (中澤 (1995)). このような観点からこれまでに様々な並列計算機が開発されてきた. ここでは並列計算機の代表的な分類法や用語, 並列処理における評価方法, 各システムの概要などについて述べる.

2.1 並列計算機の種類

今回用いた計算機は, いずれも MIMD 方式の分散メモリ型並列計算機である. MIMD (Multiple Instruction stream Multiple Data stream) とは, 独立の命令に基づいて, 各プロセッサが複数のデータに対して同時に処理を行なう方式のことである. 各プロセッサが独立に動作するが, 他のプロセッサとの途中結果の受け渡しをするために, プロセッサ同士の通信が必要になる (Flynn (1972)).

また, 分散メモリ型 (Distributed memory system) とは, 各プロセッサが独立にそれぞれメインメモリを持ったシステムのことである. これに対して, 共通のメインメモリを多数のプロセッサから直接共有出来るシステムを, 共有メモリ型 (Shared memory system) と呼ぶ.

また, 並列計算機のプロセッサ間の相互結合網による分類の方法もある. 主要なもののいくつかを図 1 に示す (中澤 (1995)).

2.2 並列処理の評価基準

1 つの仕事の分割数を P とし P 台の演算装置 (Processing Unit: 以下 PU) により分担して並列に実行したと仮定し, その時の処理時間を T_P とする. 逐次処理は $P=1$ のときに相当し, その処理時間を T_1 とすると, 理想的には並列処理時間は

$$(2.1) \quad T_P = \frac{T_1}{P}$$

のようになり, 従って速度は P 倍に増すと予測される.

しかし一般的には, そうならない場合が多い. 例えば, 1 つの仕事をも均等に P 分割できるとは限らず, 場合によっては分割した仕事 (タスク) のサイズに不均衡が生ずる. また, 分割したためにタスク間で何らかのデータの受渡しが必要になる場合がある. また, ある PU がデータ

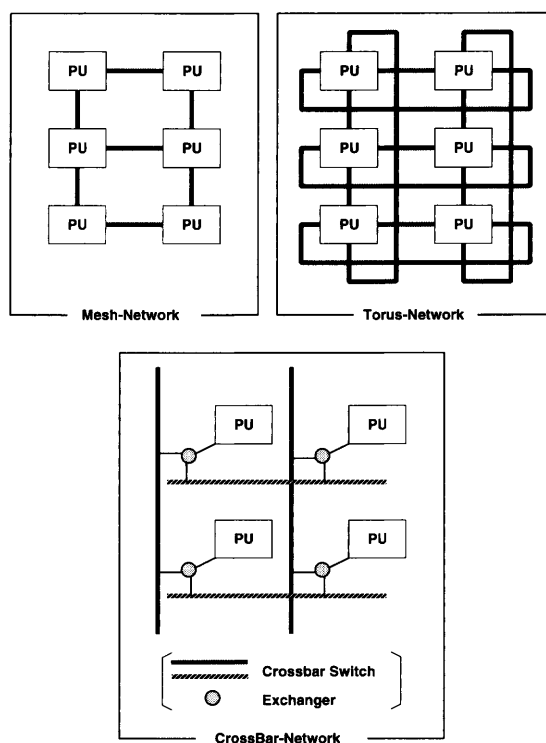


図 1. 並列計算機の主要な相互結合網。

を必要とした時に、他の PU がビジー状態にありデータがすぐに取得出来ないことがある。この時、この PU は何もせずデータを受けとるまで待ち状態になることを意味し、このような場合、効率の著しい低下を招くことがある。このようなデータ転送やそれに伴う同期などは逐次的処理には含まれていない並列処理特有の無駄な処理であり、これをオーバーヘッドという。

このオーバーヘッドのため実際には、

$$(2.2) \quad T_P \geq \frac{T_1}{P}$$

となる。そこで、並列処理効率を

$$(2.3) \quad \alpha \equiv \frac{T_1}{PT_P}$$

とする ($0 \leq \alpha \leq 1$)。また、

$$(2.4) \quad S \equiv \frac{T_1}{T_P}$$

として速度向上比を定義すると、

$$(2.5) \quad S = \alpha P$$

と表すことが可能である。

このオーバーヘッドを削減し並列計算を効率良く行なわせる、即ち並列処理効率を向上させるためには、仕事を如何に並列処理可能な部分に分割するか、与えられたデータを如何にして PU にマッピングするかなどが問題となる。

2.3 MPI

MPI は Message Passing Interface を意味し、メッセージ通信のプログラムを記述するために広く使われる標準に発展することを目的に、MPI Forum により策定された。MPI Forum には 40 以上の組織が参加しており、1992 年 11 月より、メッセージ通信の標準ライブラリセットの定義および検討がなされている。メッセージ通信の標準を確立することの主要な利点は、移植性と使い易さである (MPIF (1995))。

本研究のパソコン・クラスタで用いた MPI の環境は、Ohio Supercomputer Center で開発された LAM (Local Area Multicomputer) である。LAM はネットワーク接続されたコンピュータのための並列処理環境及び開発システムであり、拡張モニタリング、デバッグツールなどによってサポートされている。(尚、LAM は、ftp.osc.edu より入手可能である。)

2.4 CP-PACS

CP-PACS (以下 CP) は、QCD 計算などの物理計算を高速に行なうために、筑波大学と日立製作所により開発された並列計算機である。結合網は、多数のクロスバスイッチを x, y, z 各方向に配列した 3 次元ハイパクロスバ (Hyper Crossbar) である。演算処理を並列に行なう PU (Processing Unit) と、入出力を分散処理する IOU (Input/Output Unit) は各方向のクロスバを繋ぐ Exchanger に接続され、最大三段のクロスバを経由することにより任意のパターンの PU 間データ転送が可能である。各 PU 間を繋ぐ結合網の転送ピーク性能は 300 MBytes/sec である。CP の構成図を図 2 に示す。

CP の各 PU が搭載するプロセッサには、PVP-SW (Pseudo Vector Processor based on Slide Window) 機能が搭載されている。この PVP-SW 機能により、主記憶のアクセス遅延が隠蔽され、スーパースカラ・プロセッサでありながら効率のよいベクトル処理が可能となり、300 MFLOPS のピーク速度を得ている (中澤 他 (1996))。

本研究では、16 台の PU を用いた測定まで行なった。

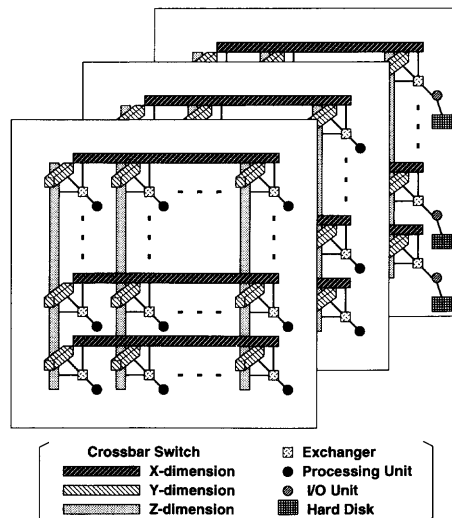


図 2. 並列計算機 CP-PACS の構成図。

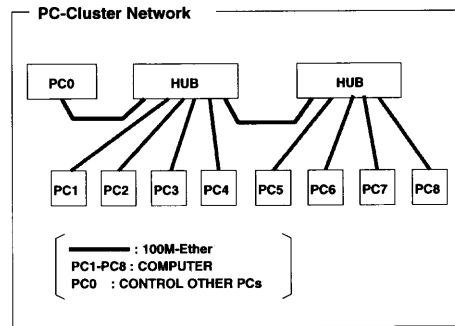


図3. パソコン・クラスタの機器構成図。

表1. 各計算機システムの仕様。

システム	CP-PACS	RS/6000 SP	パソコン・サーバ
ピーク演算スピード /1PU	300MFLOPS	266.4MFLOPS	200MFLOPS
メモリ /1PU	64MBytes	256MBytes	64MBytes
キャッシュ・メモリ /1PU	512KBytes	128KBytes	256KBytes
ピーク転送スピード	300MBytes/sec	150MBytes/sec	12.5MBytes/sec

2.5 RS/6000 SP

並列計算機 RS/6000 SP (以下 SP) は、IBM により開発された並列計算機である。各 PU は、RS/6000 のプロセッサ、メモリなどで構成されており、全体は複数の PU が格納されたフレームを接続することによって構成されている。システム全体の管理は、SP と接続されたコントロール・ワークステーションによって一元的に行なわれている。各 PU 間の通信は、SP スイッチと呼ばれるネットワークを介して行なわれるが、このピーク性能は 150 MBytes/sec である。また、プロセッサのピーク速度は 266.4 MFLOPS である。

今回は、16 台の PU を用いた測定まで行なった。

2.6 パソコン・クラスタ

パソコン・クラスタ (以下 PC) の機器構成は、AT 互換機 8 台を 100 Mbits/sec イーサネットに繋いだものである。各パソコンの MPU は、Intel PentiumPro 200 MHz、メモリは 64 MBytes である。MPU のピーク速度は 200 MFLOPS である (図3)。NFS によりホームディレクトリの共有をしており、データの入出力などが行ないやすくなっている。OS は FreeBSD 2.2.1, MPI は LAM6.0 を用いた。尚、他の 1 台のパソコンから LAM の立ちあげ、プログラムの実行などを行ない、8 台のパソコンは計算だけを行なうようにした。

各システムの仕様を表1に示す。尚、表中のキャッシュ・メモリとは、PU と主記憶の間に置かれる、高速・小容量のメモリである。よく使う主記憶内の情報の写しを一時的にキャッシュ・メモリに保存しておき、主記憶の速度不足を補うことを目的としている (中澤 (1995))。

3. 統計処理の並列化

我々はこれまでに、一般的によく行なわれている統計解析のうち、基本統計量を求める処理、重回帰分析、主成分分析について並列化を行ってきた (下平 他 (1995))。しかし、これらの

処理は並列計算機 QCDPAX (白川 (1989)) 上で専用の言語を用いて開発されたため、汎用性がなかった。そこで今回はプログラムに汎用性を持たせるため、これらの処理を MPI を用いて開発し直した。以下にアルゴリズムの概要を示す。

3.1 基本統計量の並列処理の方法

統計量のうち、最大値、最小値、平均、分散、標準偏差、変動係数、共分散、相関係数をまとめて基本統計量と呼ぶことにする。

プログラムではまず、PU[0] で各 PU 毎のデータの割り当て個数を計算し、PU[0] が各 PU へ割り当て個数分のデータを転送する。データの入力処理は PU[0] で行なう。これは、NFS などの環境がないような場合でも適用出来るような汎用性を考慮して行なった。データの割り当て方法は、全部で N 個のデータを P 台の PU に N/P 個ずつ割り当てた。そして、全 PU へのデータの入力が完了した後一斉に計算をスタートし、その結果を各 PU 毎の表として出力した。

基本統計量の処理では、各 PU で計算した値の比較・加算などのために、PU 間で $\log_2 P$ 回のデータ転送を行なっている。これが、並列処理を行なったために発生する、PU 数による計算のオーバーヘッドとなる。

3.2 重回帰分析の並列処理の方法

重回帰分析では、まず 3.1 で述べた手順で基本統計量を求め、その後 Beaton 法により解を求めた。

行列のデータは各行を各 PU に割り当てた。つまり、行列の第 0 行を PU[0] に、第 1 行を PU[1] にというように順に割り当てていき、第 $P-1$ 行を PU[$P-1$] に割り当てるまで行なう。そして、第 P 行からは再び PU[0] から順に割り当てていくという方法である。各 PU は、自分に割り当てられた行がピボット行になった場合は、その行を全 PU にブロードキャスト (複数 PU への同時通信) する。よって、データが M 変数の場合、 M 回のブロードキャストが発生し、これが重回帰分析の並列処理によるオーバーヘッドとなる。

3.3 主成分分析の並列処理の方法

主成分分析のプログラムでは、まず 3.1 に示した方法で基本統計量を求める。

そして、分散共分散行列を用いて解くか、相関行列を用いて標準化されたデータとして解くかをユーザが選択する。その行列を用いてヤコビ式を解き、得られた値から寄与率、累積寄与率、主成分得点を求める。

尚、今回の実験では、どちらの手法で解くかは予めプログラム中で与えておき測定した。

主成分分析では、解析の前に一度各 PU が受け持つデータの割当てを決めてしまえば、計算途中には通信は行なわない。よって、処理の途中でオーバーヘッドは発生しない。また、固有値を求める処理を含んでいるので、データによっては収束に時間がかかり処理時間が長くなる場合もあり得る。

4. 測定結果

前記の機器を用いて、変量数とサンプル数を変えながら実行時間を測定した。文中、PU 数を P 、データの変量数を M 、サンプル数を N と表すことにする。また、重回帰分析、主成分分析とも、その処理時間はその処理に特有な時間を表しており、各処理の前に行なう基本統計量を求める時間は含めていない。

4.1 各環境の特性

4.1.1 CP-PACS の測定結果について

基本統計量の処理では、表 2 に示すように、データが大量の場合には並列処理効率は 100% もしくは 100% 以上を示した。この測定結果のうち、PU 数と並列処理効率との関係をグラフで表したものが図 4 である。

特に、 $M=50, 100$ で $N=100000$ の時、 $P=8$ の時の効率が一番良いが、これはデータがキャッシュ・メモリに入り切るからだと思われる。 $P=16$ の時の効率が落ちているが、これはキャッシュ・メモリの利用という点では $P=8$ の時と同じ条件であるが、PU が増えた分転送回数も増えるので、結果的に効率が落ちていると思われる。

同じ PU 数で比較すると、データ量が多いほど効率は良くなっている。これは、データ量が増加しても通信回数は P によって一定だからである。また、データ量が少ない時には効率がよくないが、これは PVP 機構は、ある程度の量のデータ (配列) を計算しないと性能を発揮しないからだと思われる。

重回帰分析では、行列計算のたびにブロードキャスト処理が必要になるので、同じデータを処理する場合 PU が多いほど効率は悪化している。

$P=16$, $N=100000$ の時を比べていくと、 M が小さい方が効率は良くなっている。これは、データがキャッシュ・メモリに入り切ってしまったために計算時間が非常に短くなったので、行列計算の転送回数などが多い処理の方が遅くなってしまったためであると思われる。 $P=8$ までは M が増えるほど効率は改善されているが、これはデータがキャッシュ・メモリに入り切らず、全処理に対する計算時間の割合が多いために転送にかかる時間よりも、短縮された計算時間の影響の方が大きいからだと思われる (図 5)。

主成分分析では、計算途中には通信を行わないので、 $M=2, 10$ の時にはほぼ 100% の並列処理効率を得られた。 $M=50, 100$ と増加すると効率は悪化してしまうが、これは並列化出来ない (全 PU が行なっている) 処理の割合が並列化で高速に処理できる部分の割合を越えてしまうからだと思われる (図 6)。

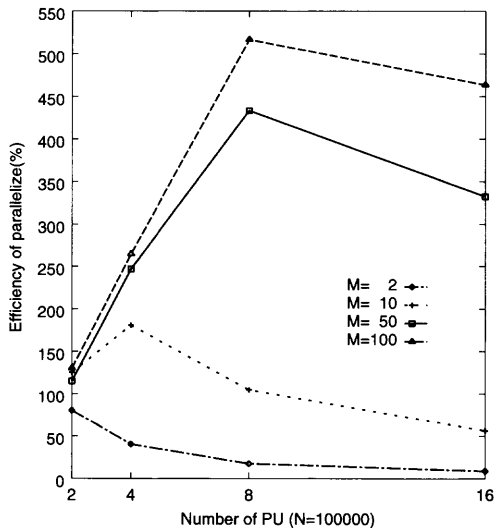


図 4. 基本統計量の並列処理効率 ($N=100000$) (CP-PACS).

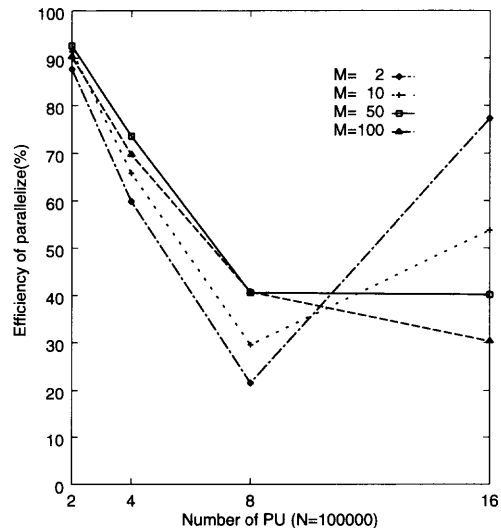
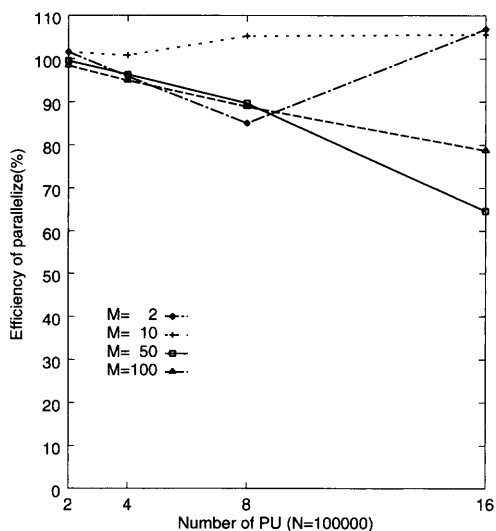
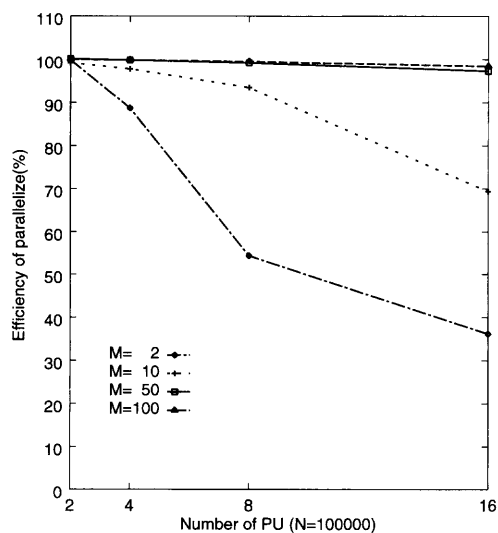


図 5. 重回帰分析の並列処理効率 ($N=100000$) (CP-PACS).

表2. 基本統計量の並列処理の速度向上比, 並列処理効率 (CP-PACS).

M		2				10			
N		1000	10000	50000	100000	1000	10000	50000	100000
PU2/PU1	speed-up	0.40	0.70	2.16	1.61	1.61	2.27	3.41	2.52
	効率 (%)	20.00	34.96	107.89	80.30	80.36	113.34	170.43	125.92
PU4/PU1	speed-up	0.25	0.47	2.05	1.61	0.63	1.55	5.94	7.23
	効率 (%)	6.25	11.65	51.24	40.15	15.63	38.64	148.48	180.82
PU8/PU1	speed-up	0.13	0.20	1.32	1.41	0.28	0.87	4.13	8.37
	効率 (%)	1.56	2.50	16.53	17.67	3.52	10.90	51.58	104.56
PU16/PU1	speed-up	0.48	0.83	4.47	1.39	0.75	2.15	5.49	8.97
	効率 (%)	2.97	5.17	27.92	8.72	4.70	13.43	34.29	56.05

M		50				100			
N		1000	10000	50000	100000	1000	10000	50000	100000
PU2/PU1	speed-up	1.29	1.79	3.86	2.30	1.52	1.99	4.23	2.61
	効率 (%)	64.59	89.49	193.12	115.10	76.19	99.50	211.71	130.59
PU4/PU1	speed-up	0.94	2.42	14.42	9.89	1.48	3.38	17.39	10.58
	効率 (%)	23.49	60.58	360.58	247.30	36.92	84.47	434.76	264.49
PU8/PU1	speed-up	0.53	2.44	22.08	34.65	1.07	5.11	32.05	41.33
	効率 (%)	6.68	30.52	276.03	433.16	13.33	63.90	400.64	516.58
PU16/PU1	speed-up	0.59	2.51	29.66	53.15	0.96	5.57	53.30	74.17
	効率 (%)	3.68	15.69	185.38	332.22	5.99	34.78	333.13	463.55

図6. 主成分分析の並列処理効率 ($N=100000$) (CP-PACS).図7. 基本統計量の並列処理効率 ($N=100000$) (RS/6000 SP).

4.1.2 RS/6000 SP の測定結果について

基本統計量の処理では, データ量が多い場合には並列処理効率は, ほぼ 100%を示した. 特に,

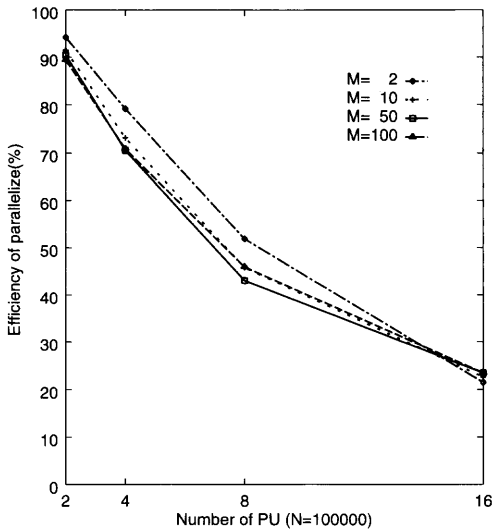


図 8. 重回帰分析の並列処理効率 ($N=100000$) (RS/6000 SP).

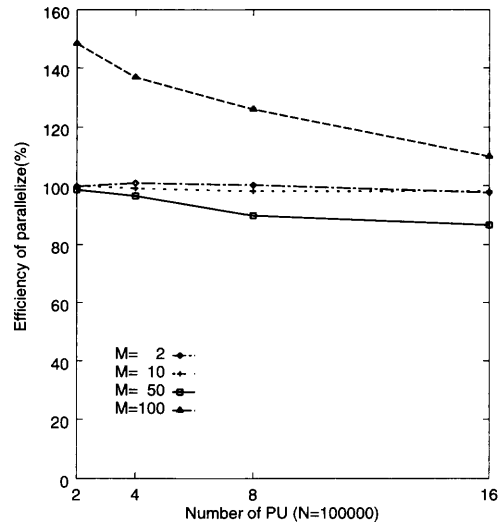


図 9. 主成分分析の並列処理効率 ($N=100000$) (RS/6000 SP).

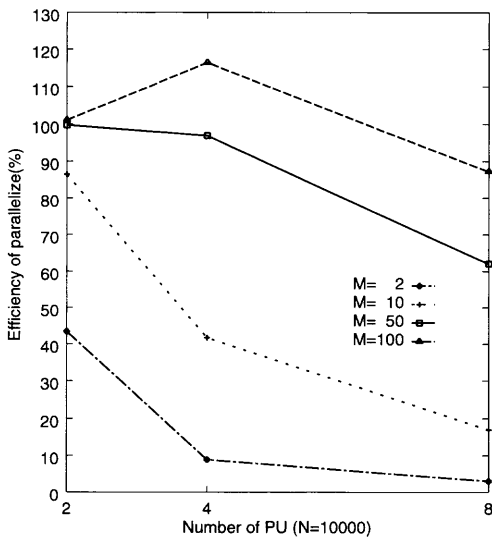


図 10. 基本統計量の並列処理効率 ($N=10000$) (パソコン・クラスタ).

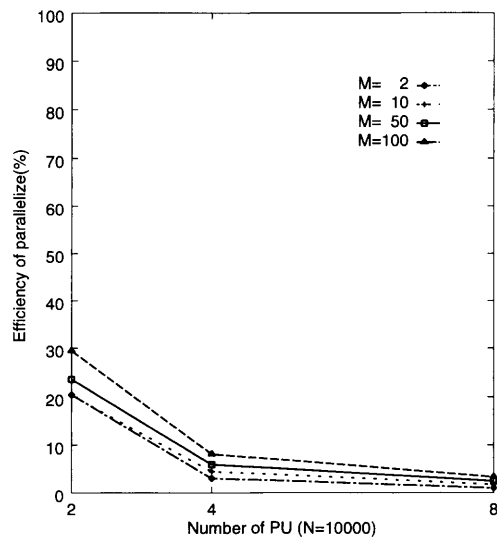


図 11. 重回帰分析の並列処理効率 ($N=10000$) (パソコン・クラスタ).

$M=100, N=100000$ の時には、ほぼ 99% の効率を得ている。データ量が少ない時には、 $P=4$ をピークに効率は悪化しているが、これは計算量が少ないために通信の影響が大きく出てくるからである。 $M=100, N=100000$ の時でも、PU がもっと増えれば効率自体は悪化すると思われる (図 7)。

重回帰分析は、CP とよく似た傾向を示した (図 8)。

主成分分析も、ほぼ CP と同じような傾向が見られる。 $M=100, N=100000$ の時、 $P=16$ で

100%以上の効率を示しているが、これもPUが増えるにしたがって悪化するとと思われる(図9)。

4.1.3 パソコン・クラスタの測定結果について

基本統計量では、データ量が少ない時はPUが増えると効率は急速に悪化する。これは、計算量(時間)が少なく通信が遅いためである。キャッシュ・メモリの容量から考えて、一番効果があるのは $P=8$ の時だと思われるが、通信速度が非常に遅いので、その兼ね合いで $P=2$ もしくは $P=4$ の時の効率の方が良い(図10)。

重回帰分析では、行列計算でブロードキャスト処理が頻繁に発生するので、非常に効率は悪い。大量のデータを $P=2$ で処理した場合でも、20~30%前後である(図11)。

主成分分析では、 $M=10, N=10000$ では、ほぼ100%の効率を得ているが、 $M=50, 100$ と増加すると効率は71.6%, 54.8%と悪化していく。これも、他のマシンでの傾向と同じように、並列化出来ていない M に依存する処理の割合が大きくなっていくためである。

よって、 M が増加するほど効率は悪化し、 N が増加するほど効率は向上する(図12)。

尚、パソコンクラスタでは、メモリ容量の関係で1台当たり約10000サンプルまでの処理しが行なうことが出来なかった。従って、効率算出にはその測定結果を用いた。

4.2 CP-PACS, RS/6000 SP, パソコン・クラスタの相対性能評価

ここでは、各処理における各測定環境ごとの実時間の比較とその考察を行なう。(尚、表中の「***」の記号は、パソコン・クラスタにおいて、PU台数やメモリ容量の関係上測定出来なかったことを示す。)

4.2.1 基本統計量

表3より、 $M=2, 10$ と小さい時は、 $P=4, 8$ でSPの方がCPよりも速い。また、 $M=50, 100$ と大きくなっていくとCPの方が大幅に速くなっていく。これは、データ量が少ない時はCPのPVP機能は効果が発揮されず、大きなデータを扱う場合に効果が得られるからである。また、

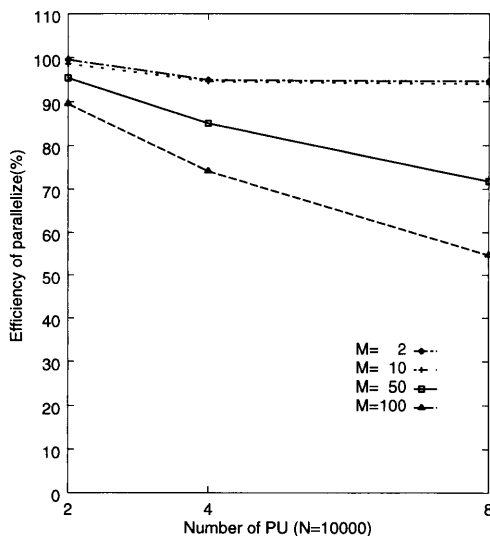


図12. 主成分分析の並列処理効率 ($N=10000$) (パソコン・クラスタ)。

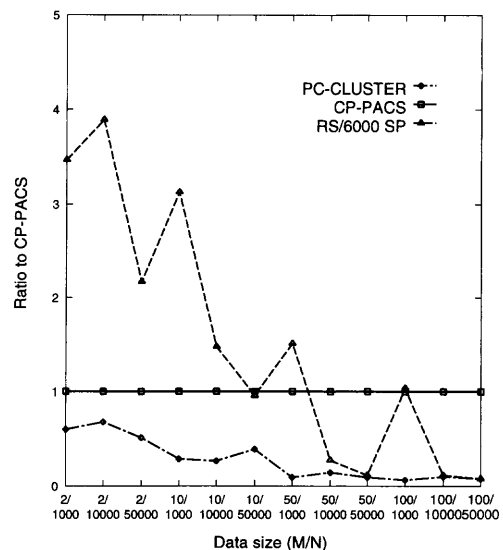


図13. 各環境における基本統計量の処理速度の比 ($P=8$)。

表 3. CP-PACS, RS/6000 SP, PC-CLUSTER の速度比 (基本統計量).

M	N	PU	CP	SP	PC	M	N	PU	CP	SP	PC
2	1000	1	1	1.465	5.290	50	1000	1	1	0.091	0.110
		2	1	1.963	1.370			2	1	0.191	0.115
		4	1	2.921	0.560			4	1	0.623	0.089
		8	1	3.467	0.700			8	1	1.512	0.111
		16	1	0.127	***			16	1	0.676	***
	10000	1	1	0.455	0.629		10000	1	1	0.082	0.074
		2	1	0.858	0.781			2	1	0.091	0.082
		4	1	1.840	0.481			4	1	0.140	0.118
		8	1	3.892	0.749			8	1	0.272	0.150
		16	1	0.302	***			16	1	0.535	***
	50000	1	1	0.797	***		50000	1	1	0.332	***
		2	1	0.636	0.421			2	1	0.171	0.154
		4	1	1.124	0.424			4	1	0.091	0.078
		8	1	2.170	0.583			8	1	0.117	0.092
		16	1	0.463	***			16	1	0.170	***
	100000	1	1	0.959	***		100000	1	1	0.425	***
		2	1	1.193	***			2	1	0.370	***
		4	1	2.117	***			4	1	0.172	***
		8	1	2.951	***			8	1	0.097	***
		16	1	3.988	***			16	1	0.125	***
10	1000	1	1	0.963	1.00	100	1000	1	1	0.065	0.085
		2	1	0.687	0.295			2	1	0.141	0.088
		4	1	2.007	0.232			4	1	0.443	0.075
		8	1	3.129	0.335			8	1	1.030	0.080
		16	1	0.577	***			16	1	0.798	***
	10000	1	1	0.235	0.188		10000	1	1	0.074	0.070
		2	1	0.194	0.143			2	1	0.074	0.071
		4	1	0.557	0.202			4	1	0.091	0.096
		8	1	1.477	0.290			8	1	0.119	0.095
		16	1	0.538	***			16	1	0.274	***
	50000	1	1	0.559	***		50000	1	1	0.333	***
		2	1	0.321	0.231			2	1	0.157	0.134
		4	1	0.357	0.205			4	1	0.077	0.071
		8	1	0.957	0.408			8	1	0.082	0.071
		16	1	1.080	***			16	1	0.100	***
	100000	1	1	0.690	***		100000	1	1	0.411	***
		2	1	0.543	***			2	1	0.315	***
		4	1	0.373	***			4	1	0.155	***
		8	1	0.616	***			8	1	0.079	***
		16	1	0.853	***			16	1	0.087	***

PUが増えると、ハイパクロスバ構造により指定したPUに一回でデータを転送できるCPの方が有利になる。

各環境での速度比は、 $P=8, M=100, N=50000$ の時、 $CP:SP:PC=1:0.082:0.071$ 。また、 $P=16, M=100, N=100000$ の時、 $CP:SP=1:0.087$ 。キャッシュ・メモリやCPのPVP機能などの影響があり、データ量やPU数によって完全に同じ割合にはなっていない (図 13)。

4.2.2 重回帰分析

表4より、 M, N がいくつでも、 $P=8$ まではSPの方が速い。ただし、データ量が多くなるほどその差は縮まり、 $M=100$ ではほぼ同じになる。この内容を詳しく見ると、通信を頻繁に行なう部分で差がついてしまっている。本来、ハイパクロスバ構造を取るCPの方が高速なはずであるが、MPIのコンパイラの最適化が不十分なのか、少ないデータを頻繁に送るという処理では高速性が発揮されていない。データ量が増加すると差が縮まっているが、これは計算部分の実

表4. CP-PACS, RS/6000 SP, PC-CLUSTERの速度比(重回帰分析)。

M	N	PU	CP	SP	PC	M	N	PU	CP	SP	PC
2	1000	1	1	33.523	42.005	50	1000	1	1	1.726	1.034
		2	1	22.026	1.958			2	1	1.384	0.155
		4	1	11.647	0.803			4	1	1.582	0.106
		8	1	17.836	1.097			8	1	1.954	0.139
		16	1	0.867	***			16	1	0.772	***
	10000	1	1	5.882	4.051		10000	1	1	1.460	0.491
		2	1	8.104	1.650			2	1	1.396	0.193
		4	1	11.720	0.799			4	1	1.543	0.116
		8	1	14.856	1.019			8	1	1.850	0.136
		16	1	0.855	***			16	1	0.792	***
	50000	1	1	3.824	***		50000	1	1	1.513	***
		2	1	3.909	1.597			2	1	1.439	0.272
		4	1	7.090	0.913			4	1	1.477	0.154
		8	1	10.795	0.969			8	1	1.710	0.153
		16	1	1.013	***			16	1	0.849	***
	100000	1	1	3.524	***		100000	1	1	1.535	***
		2	1	3.783	***			2	1	1.498	***
		4	1	4.664	***			4	1	1.473	***
		8	1	8.526	***			8	1	1.631	***
		16	1	0.976	***			16	1	0.898	***
10	1000	1	1	12.066	7.410	100	1000	1	1	1.060	0.730
		2	1	5.729	0.476			2	1	1.042	0.188
		4	1	4.954	0.275			4	1	1.065	0.094
		8	1	6.107	0.391			8	1	1.261	0.104
		16	1	0.794	***			16	1	0.738	***
	10000	1	1	4.009	1.372		10000	1	1	1.002	0.448
		2	1	3.220	0.412			2	1	1.010	0.191
		4	1	4.490	0.299			4	1	1.041	0.099
		8	1	5.840	0.396			8	1	1.202	0.103
		16	1	0.787	***			16	1	0.735	***
	50000	1	1	2.493	***		50000	1	1	0.957	***
		2	1	2.561	0.544			2	1	0.943	0.199
		4	1	3.264	0.340			4	1	0.980	0.116
		8	1	4.289	0.378			8	1	1.125	0.109
		16	1	0.918	***			16	1	0.752	***
	100000	1	1	2.442	***		100000	1	1	0.957	***
		2	1	2.440	***			2	1	0.948	***
		4	1	2.714	***			4	1	0.973	***
		8	1	3.767	***			8	1	1.077	***
		16	1	1.026	***			16	1	0.738	***

行時間の割合が大きくなっていくので、PVP 機能で高速な処理が出来る分だけ CP の処理効率が向上しているからだと思われる。

$P=16$ の時 SP の速度が急に遅くなっているが、これは SP の構造上の特徴によると思われる。SP は 16PU を一つのフレームに格納し、それを繋ぎ合わせるという構造を取っている。各フレーム内には、ゲートウェイや単体演算用の PU が設定されており、全部を並列に使用できる訳ではない。そこで、16PU を使用した計算をする場合は、二つ以上のフレームをまたいで処理

表5. CP-PACS, RS/6000 SP, PC-CLUSTER の速度比 (主成分分析)。

M	N	PU	CP	SP	PC	M	N	PU	CP	SP	PC
2	1000	1	1	2.878	7.025	50	1000	1	1	1.471	0.434
		2	1	4.580	9.742			2	1	1.569	0.523
		4	1	10.404	18.382			4	1	1.461	0.674
		8	1	41.296	65.917			8	1	1.419	0.652
		16	1	2.298	***			16	1	1.643	***
	10000	1	1	1.971	4.338		10000	1	1	1.339	0.309
		2	1	2.195	4.836			2	1	1.388	0.352
		4	1	2.834	6.260			4	1	1.326	0.349
		8	1	5.435	11.821			8	1	1.447	0.424
		16	1	1.777	***			16	1	1.570	***
	50000	1	1	2.005	***		50000	1	1	1.253	***
		2	1	1.955	2.216			2	1	1.260	0.142
		4	1	2.080	4.452			4	1	1.269	0.296
		8	1	2.686	4.800			8	1	1.373	0.332
		16	1	1.829	***			16	1	1.545	***
	100000	1	1	1.988	***		100000	1	1	1.280	***
		2	1	1.954	***			2	1	1.270	***
		4	1	2.095	***			4	1	1.282	***
		8	1	2.349	***			8	1	1.286	***
		16	1	1.820	***			16	1	1.717	***
10	1000	1	1	2.306	1.511	100	1000	1	1	1.109	0.358
		2	1	3.615	2.466			2	1	1.093	0.375
		4	1	8.687	5.811			4	1	1.063	0.388
		8	1	21.687	15.145			8	1	0.836	0.395
		16	1	1.282	***			16	1	0.909	***
	10000	1	1	1.951	1.226		10000	1	1	1.340	0.306
		2	1	2.224	1.399			2	1	1.289	0.317
		4	1	2.709	1.679			4	1	1.128	0.333
		8	1	5.372	3.470			8	1	1.039	0.350
		16	1	1.799	***			16	1	1.096	***
	50000	1	1	1.913	***		50000	1	1	1.399	***
		2	1	1.892	1.095			2	1	1.376	0.287
		4	1	2.047	1.278			4	1	1.322	0.301
		8	1	2.323	1.360			8	1	1.244	0.312
		16	1	1.788	***			16	1	1.213	***
	100000	1	1	1.959	***		100000	1	1	0.750	***
		2	1	1.930	***			2	1	1.413	***
		4	1	1.927	***			4	1	1.349	***
		8	1	1.829	***			8	1	1.325	***
		16	1	1.823	***			16	1	1.309	***

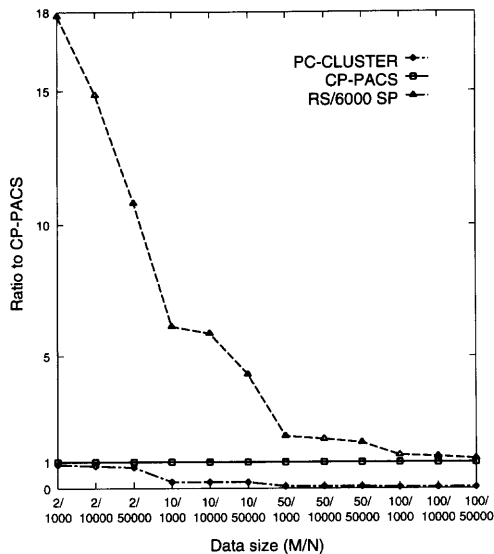


図 14. 各環境における重回帰分析の処理速度の比 ($P=8$).

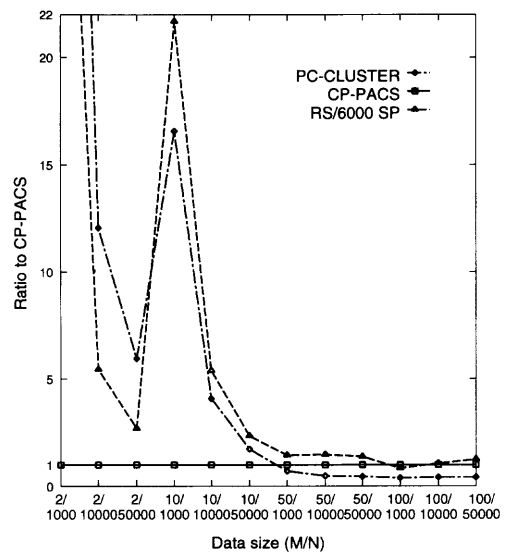


図 15. 各環境における主成分分析の処理速度の比 ($P=8$).

を行なうことになり、その分だけオーバーヘッドが大きくなっているのではないと思われる。

PCでの効率が極端に悪くなっているが、これは行列計算で頻繁にブロードキャスト処理が行なわれるためと思われる。PCで使用したLAM/MPIシステムでは、一般的なTCP/IP通信を用いて通信を行なっているため、ブロードキャスト処理とは言っても実際は台数分だけ一対一通信を繰り返しているだけである。従って、PCの台数が増えるほどブロードキャスト処理にかかる時間は長くなっていき、全体の効率は悪化していくのである。

各環境での速度比は、 $P=8$, $M=100$, $N=50000$ の時、 $CP:SP:PC=1:1.125:0.109$ である。また、 $P=16$, $M=100$, $N=100000$ の時、 $CP:SP=1:0.738$ である。重回帰分析の処理も、システムの構造的な特性からデータ量によって同じ割合にはなっていない(図14)。

4.2.3 主成分分析

表5より、全体的にSPの方が高速である。特に、データ量が少ない時の方がより高速である。また、データ量が多くても、約1.2~1.3倍前後の差がある。これは、CPのコンパイラではコンパイル時にPVP機能を使えるかどうかの判別をしているが、まだ完全ではなく、ループ処理でPVP機能を使えていない場合が見られることと、PVP機能を使える場合でも、少ないデータでは効果が発揮できないからである。ただし、今回実験を行なったよりも更に大量のデータを処理していく場合には、PVP機能を使用出来ている部分の処理の高速化が望めるので、差は縮まっていくと思われる。

また、重回帰分析の時に見られたような、SPのPUが16台の時の顕著な速度低下は見られない。これは、主成分分析の処理の中では、途中に通信処理が含まれないからである。

各環境での速度比は、 $P=8$, $M=100$, $N=50000$ の時、 $CP:SP:PC=1:1.244:0.312$ である。また、 $P=16$, $M=100$, $N=100000$ の時、 $CP:SP=1:1.309$ である。計算途中に通信処理がないという特性のため、PCでもそれほど大きな速度低下はなく、実用に耐え得る速度が得られた(図15)。

5. 結 論

基本統計量では、データの通信回数は決まっている。データ量が少ない時、特に $M=2, 10$ など変数数が小さい時には PC でも実用に耐え得る速度で処理できる。しかし、 M, N が大きくなってくるとその差は大きくなり、 $M=100, N=100000$ 前後では PVP 機能により高速なベクトル演算が出来、ハイパクロスバ構造でデータ転送の高速な CP が一番有利である。

重回帰分析では、ブロードキャスト処理によりデータを頻繁にやりとりする。そのため、ネットワーク速度の遅い PC では非常に時間がかかる。また、本来一番高速な通信能力を持っているはずの CP は、ユーザによるチューニングが行なわれることをある程度前提として開発されているのか、今回のように MPI で記述したプログラムではスペック通りの速度を発揮できなかった。そのため、今回測定した範囲内では SP が一番適しているという結果が出た。

主成分分析では、基本的に処理の途中には通信を行っていない。よって、計算の速度のみが影響すると思われるが、今回コンパイラの関係で CP で PVP 機能を使い切れていない部分があり、SP の方が高速であるという結果が出た。また、通信がないおかげで PC の処理の速度低下もそれほど大きくはなく、実用に耐え得るレベルの速度が得られた。

今回、プログラムの汎用的な使用を考え通信ライブラリとして MPI を使用し、敢えて各機種固有のライブラリを用いたチューニングなどは行なわなかった。その影響で、CP では十分に能力を発揮できなかった部分もあるが、様々なプラットフォームでそのまま利用できるプログラムということの利点は大きいと考える。

謝 辞

本研究は、統計数理研究所共同研究 (9-共研 A-122) 「統計データ解析の並列処理」に基づくものである。本研究で用いた並列計算機 CP-PACS は、文部省科学研究費補助金 08NP0101 により開発された。

参 考 文 献

- Agrawal, R., Imielinski, T. and Swami, A. (1993). Mining association rules between sets of items in large databases, *Proceedings of ACM SIGMOD International Conference on Management of Data*, Washington D.C.
- Flynn, M. J. (1972). Some computer organization and their effectiveness, *IEEE Trans. Comput.*, **C-21**(9), 948-960.
- MPIF (Message Passing Interface Forum) (1995). MPI: A Message-Passing Interface Standard, <http://www.mpi-forum.org/>.
- 村上征勝, 田村義保 (1998). 『パソコンによるデータ解析』, 朝倉書店, 東京.
- 中澤喜三郎 (1995). 『計算機アーキテクチャと構成方式』, 朝倉書店, 東京.
- 中澤喜三郎 (1996). 計算物理学研究用並列計算機: CP-PACS, 応用数理, **6**, 17-28.
- 中澤喜三郎, 中村 宏, 朴 泰祐 (1996). 超並列計算機 CP-PACS のアーキテクチャ, 情報処理, **37**, 18-28.
- 坂元慶行 (1985). 『カテゴリカルデータのモデル分析』, 共立出版, 東京.
- 下平文彦, 小林 覚, 白川友紀, 田村義保 (1995). 統計データ解析の並列処理, 計算機統計学, **8**, 15-26.
- 白川友紀 (1989). QCDPAX のアーキテクチャ, ソフトウェア, 応用情報学研究年報, **14**(2), 127-139.
- 杉山高一 (1983). 『多変量データ解析入門』, 朝倉書店, 東京.
- 田中豊, 垂水共之, 脇本和昌 (1984). 『パソコン統計解析ハンドブック II 多変量解析編』, 共立出版, 東京.
- 渡部 洋, 鈴木規夫, 山田文康, 大塚雄作 (1985). 『探索的データ解析入門 —— データの構造を探る ——』, 朝倉書店, 東京.
- 柳井晴夫, 高木廣文 (1981). 『探索的データ解析の方法』, 朝倉書店, 東京.

Performance Evaluation on Parallel Computer Systems in Executing Statistical Data Analysis

Fumihiko Shimodaira

(Department of Engineering, University of Tsukuba)

Tomonori Shirakawa

(Institute of Engineering Mechanics, University of Tsukuba)

Yoshiyasu Tamura

(The Institute of Statistical Mathematics)

In recent years, in step with the increase in computing resources, data sets have grown in size and complexity in many situations. The analysis of those data is becoming very important. It will take very long time to analyze mass data precisely.

We aim at high speed execution of statistical data analysis on several parallel computers. The performances of a personal computer cluster, the parallel computer "CP-PACS" and "RS/6000 SP" were compared by using our statistical data analysis programs.

We evaluate the efficiency of parallel processing of basic statistics, multiple regression analysis and principal component analysis.

As the result, the efficiencies of parallel processing were over 100% (463.6%), 30.4% and 78.7% respectively on 16 PUs of the CP-PACS with 100 variables and 100000 samples. The efficiencies of parallel processing were 98.4%, 23.4% and over 100% (137.3%) respectively on 16 PUs of the RS/6000 SP with 100 variables and 100000 samples. The efficiencies of parallel processing were 87.1%, 3.4% and 54.8% respectively on 8 PUs of a personal computer cluster with 100 variables and 10000 samples.

Comparing each execution time, we conclude CP-PACS is fit for basic statistics analysis, RS/6000 SP is fit for multiple regression analysis which uses broadcasting data execution frequently, and RS/6000 SP is fit for principal component analysis. As principal component analysis has no communication execution, a personal computer cluster is of practical use.