

キャッシュ・オンリ・メモリ・アーキテクチャ向き ディレクトリ機構の提案

福田 宗弘[†], 村田 浩樹^{††}, 清水 茂則^{††}

[†]Department of Information and Computer Science
University of California, Irvine

^{††}日本アイ・ビー・エム（株） 東京基礎研究所

キャッシュ・オンリ・メモリ・アーキテクチャ（以下COMA）は分散共有メモリ・システムの一つで、各ノードのローカル・メモリを一種のレベル3キャッシュとして用いることで、大きな共有メモリ領域を提供しようとするものである。しかし一般に階層ディレクトリ機構を用いるためメモリ遅延が長くなるという欠点を持つ。さらにメモリ空間をキャッシュとして使いきってしまった場合、新たなメモリ空間を得るためにOSのサポートが必要となる。本論文では、COMAのメモリ・アクセス遅延を減少させるための非階層ディレクトリ機構と、OSのサポートを不要とするためのデータ割付プロトコルの2つを提案する。

Flat Directory Scheme for Cache-Only Memory Architecture

Munehiro Fukuda[†], Hiroki Murata^{††}, and Shigenori Shimizu^{††}

[†]Department of Information and Computer Science
University of California, Irvine
Irvine, CA 92717, U.S.A.

^{††}Tokyo Research Laboratory, IBM Japan Ltd.
1623-14, Shimotsuruma, Yamato-shi, Kanagawa 242, Japan
e-mail: murata@trl.ibm.co.jp

Cache-only memory architecture (COMA), a distributed shared-memory system, uses each node's local memory as a kind of third-level cache to provide a large capacity for shared data. However it has the drawback of long memory access latency because of its hierarchical directory scheme. In addition, the operating systems (OSs) must be involved in directory tag invalidations upon reclaiming a new caching area. We aim to reduce the memory access latencies without increasing the directory size or the number of inter-node messages, and to reduce OS involvement. For these purposes, we propose a right-sized flat directory scheme and a caching-space allocation protocol.

1 はじめに

密結合型マルチプロセッサ・システムは、プロセッサ間で同一メモリ空間を共有しており、データ転送が暗黙のうちになされるため並列処理に広く用いられている。しかし、それらはプロセッサ間接続がネックとなりやすく、これを避けるために複雑で高価なネットワークを必要とする場合がある。また新しいCPUが作られるたびに、設計をし直さなくてはならないことが多いという欠点もある。

近年は分散共有メモリ・システム(DSM)[5]が注目されている。DSMはネットワーク接続されたワークステーション・クラスタ上に共有メモリ・モデルを構成するもので、ハードウェア設計に大きな負担をかけることなく、並列度を向上させることができる。Fiber Channel Standard(FCS)[7]やScalable Coherent Interface(SCI)[1]などの高速な通信媒体の標準化、商業化によって、ノード間のキャッシュ・コヒーレンスをハードウェアによって維持することが実用的になりつつある。

ここで2種のハードウェア指向DSMシステムについて述べる。1つはcache-coherent nonuniform memory architecture (cache-coherent NUMA)を用いたもの、他方はcache-only memory architecture (COMA)を用いたものである。これらは一般的に分散ディレクトリ機構を用いることで、キャッシュ・コヒーレンス維持のためのメッセージをデータを共有するノードのみに送る。これによってブロード・キャストした場合よりもメッセージ数を減らしている。これらのアーキテクチャの主な違いは、cache-coherent NUMAが各ノードのローカル・メモリを通常の主記憶として利用するのにに対し、COMAはキャッシュとして用いる点である。これらの得失についてはP. Stenströmらによって、文献[6]で解析的に述べられている。

Cache-coherent NUMAを用いた例としては、Stanford's DASH マルチ・プロセッサ[3]がある。Cache-coherent NUMAシステムでは、各ノードに非階層ディレクトリと他のノードからデータをキャッシュするための特別のローカル・キャッシュが設けられる。非階層ディレクトリは自ノードのデータをキャッシュしているノードを指す。このためcache-coherent NUMAシステムはデータ検索が容易で早く、特定のネットワーク・トポロジに依存

しない。しかし他ノードからのデータのキャッシュ量がローカル・キャッシュの大きさに制限される。COMAを用いた例にはSwedish Institute of Computer Science's Data Diffusion Machine (DDM)[2]がある。COMAを採用した計算機では、データは物理アドレスとは関係なく配置され、ノード間を移動する。システムのどこかに存在する必要なデータを見つけ出すために、各ノードは、まず少数でクラスタとしてネットワーク接続される。各クラスタはメモリ・アクセスがクラスタ内で済むか済まないかを判別するディレクトリを介して上位のネットワークに接続され、規模の大きなクラスタを構成する。その大クラスタがさらにディレクトリを介して上位のネットワークへ、というように階層化して接続される。したがって、COMAはcache-coherent NUMAとは逆に、メモリ・アクセス遅延は長い、キャッシュできるデータ量が大きいという利点を持つ。

Flat COMAは両アーキテクチャの利点、すなわち短いメモリ・アクセス遅延と大容量キャッシュをを合わせ持つものとしてP. Stenströmによって文献[6]で提案されたが、実現する上で重大な問題がある。Flat COMAはCOMAの階層ディレクトリをNUMA型の非階層ディレクトリで置き換えるが、その実現方法である。単純な実現方法としてはK. Liが文献[4]で提案したfixed distributed manager algorithmによるものがある。このアルゴリズムでは、各ノードは、DSM空間中の予め割り当てられた範囲のデータについて、それらを所有するノード(オーナー・ノード)を記録するowner management tableを持つ。このためデータをアクセスする場合、そのデータのオーナー・ノードを記録しているノードを介してアクセスするため、直接アクセスできる場合に比べメッセージが1つ増え遅延が長くなる。同様にK. Liが提案したdynamic distributed manager algorithmによるものでは、各ノードが全てのDSM空間のデータのオーナー・ノードをディレクトリに記録するため直接オーナー・ノードに到達できる。しかしこのアルゴリズムには2つの欠点がある。1つは全てのノードが同時にディレクトリを更新するわけではないので、実際のオーナーがすでに変わっているのに、メッセージがフォワードされる場合がある。もう1つはノード数が増加した場合にディレクトリが巨大なものになるという点である。適正な規模のディレクトリで、メッセージ数が少なくアクセス遅

延を短くできるアルゴリズムが望ましい。もう一つの COMA の重大な問題点は、空きメモリがなくなった場合に、OS が使用されないと考えられるデータをディスクに書き出して、メモリに空きを作らなくてはならない点である。ディレクトリなどをハードウェアで構成した場合、これは OS のメモリ・マネージャ部分にハードウェア依存のコードを増加させることになり、これも少ない方が望ましい。

本論文では、flat COMA のための、適正な規模のディレクトリを実現するオーナーシップ管理機構を提案する。また空きメモリの欠乏を防ぐキャッシュ・スペース割当プロトコルについて考察を試みる。オーナーシップ管理機構は fixed distributed manager algorithm と dynamic distributed manager algorithm を統合したもので、メッセージ数とメモリ・アクセス遅延を減少させる。各ノードは2つのテーブル:dynamic directory table と fixed owner table を持つ。前者は、自ノードに持っているデータの coherence 維持と、オーナーの記録に使われる。後者は自ノードに予め割り当てられた DSM 空間のデータのオーナーを記録する。キャッシュ・スペース割当プロトコルについては、自ノードのローカル・メモリに置かれる共有データの比率をハードウェアで調節することで空きメモリの欠乏を防ぐ方法について考察する。

ここで以後の議論に備えて、語句を定義する。

リクエスト: データを要求するノード

マネージャ: DSM 空間の予め指定された範囲のデータのオーナーシップ情報を記録するノード

オーナー: データのオーナーシップを持つノード

シェアラ: データのオーナーシップは持たないがデータを共有しているノード

ネイバ: ノード ID が自ノードのものより1大きいノード

データ要求メッセージ: リクエストからオーナーへのリモート・アクセス

データ供給メッセージ: オーナからリクエストへにデータを返す

確認メッセージ: 新オーナーからマネージャへのオーナーシップ更新情報

無効化メッセージ: オーナからシェアラへの共有データ無効化

要求遅延: ローカル・プロセッサの要求から確認メッ

セージまでの時間

2章では、本論文で考える flat COMA の構成について述べる。3章では、flat COMA のための、適正な規模のディレクトリを実現するオーナーシップ管理機構について述べる。4章では、3章で述べたオーナーシップ管理機構の性能について簡単な評価をおこなう。5章では空きメモリの欠乏を防ぐキャッシュ・エリア割当プロトコルについて考察を試み、6章で簡単なまとめを述べる。

2 Flat COMA システムの構成

我々の想定する flat COMA システムは、図1に示すように、各ワークステーションが通信ボードに非階層ディレクトリを持ち適当なネットワークに接続される。ワークステーションはローカル・メモリをレベル3・キャッシュとして扱うので、このレベル3・キャッシュがデータをリプレイスする先は存在しない。あえてリプレイスするとすれば他のノードとなる。プロセッサ/キャッシュがデータをローカル・メモリに要求すると、通信ボードはそのデータがローカル・メモリに存在するかをチェックし、存在すればローカル・メモリからデータを供給する。データがローカル・メモリに存在しなかった場合、通信ボードは、そのデータの最新のオーナーへデータ要求メッセージを発行し、データ供給メッセージを待ち受ける。リモート・アクセス中のデッド・ロックを避けるために、各ノードのシステム・バスはスプリット・バス転送機能を持つ。通信ボードの主要構成要素は dynamic directory table, fixed owner table, directory controller, link control chip である。Dynamic directory table の持つ機能は、DSM アドレスから物理アドレスへの変換と、ローカル・メモリにあるデータの coherence/オーナーシップの維持である。Fixed owner table は、予め割り当てられた範囲の DSM 空間のデータのオーナーを記憶する。Directory controller は、これら2つのテーブルを更新/維持する。Link control chip は、ネットワークとのメッセージのやり取りを行なう。プロセッサ/キャッシュが要求したデータがローカル・メモリに存在しない場合、データ要求メッセージが link control chip を介してネットワークに送出される。Dynamic directory table に要求されたデータの有効なエントリがあれば、そこに記録されているオーナーにメッ

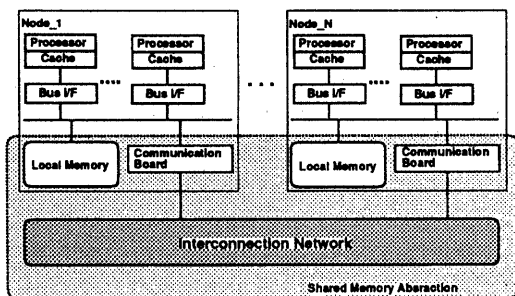


図 1: システム構成

セージが直接送られる。有効なエントリがなければ、データの fixed owner table を持つマネージャ経由で送られる。このような仕組みで、階層ディレクトリを置き換えることができる。

3 オーナシップ管理機構

COMA は、ソフトウェア・サポートによってノード間でページを共有する virtual shared memory (VSM) と、データの位置を物理アドレスから分離するという点で非常に良く似ている。したがって、VSM システムでオーナに捜し出すために提案された dynamic distributed manager algorithm は、理論的には COMA に適用することができる。このアルゴリズムでは、各ノードのページ・テーブルが、すべての DSM データのオーナシップとコヒーレンシについての情報を持っている。ノードはページ・テーブルのオーナシップ・フィールドを参照して、直接オーナにデータ要求メッセージを発行する。該当データがすでに他のノードに転送されていた場合、データ要求メッセージもそのノードへ転送される。このメッセージ転送の繰り返しは、転送に関わるノード全てにリクエストの ID を知らせることで有限回に制限できる。VSM と COMA の違いは、DSM 空間をアクセスする際に用いるアドレスの違いにある。VSM は仮想アドレスを用いて DSM 空間をアクセスするが、COMA は物理アドレスを用いる。このため VSM と COMA では、データ・アクセスの大きさが異なる。VSM は、データをページ単位で扱い、ディレクトリとしてページ・テーブルを用いるが、COMA は、データをキャッシュ・ライン単位で扱い、ディレクトリ・テーブルを用いて制御する。Dynamic

distributed manager algorithm を COMA に適用すると、VSM と比較して非常に大きなディレクトリ・テーブルが必要となる。例えば、32 ノードからなる COMA システムで、各ノードに 64Mbyte のローカル・メモリを持ち、キャッシュ・ライン長が 64byte で、ディレクトリの各エントリに 2byte 必要だとすると、ディレクトリ・テーブルとして 64Mbyte 必要になる。

我々は各ノードが、dynamic distributed manager algorithm のように全 DSM 空間に対応するディレクトリ・テーブルではなく、ローカル・メモリ空間に対応するディレクトリ・テーブルを持つべきであると考え。この場合、ディレクトリ・サイズはシステムの拡張とは関係なく、ローカル・メモリ・サイズによって決まる。前述の条件で、その大きさは 2Mbyte となり、システム中のメモリ量によらず一定である。我々はこのディレクトリを dynamic directory table と呼ぶ。このテーブルの構造は、表 1 に示すようにローカル・メモリに関係付けられた一種のフル・アソシアティブ・タグ・メモリである。データが一度ローカル・メモリに読み込まれると、dynamic directory table は、そのデータの物理アドレス、コヒーレンシ情報、と最新のオーナシップを記録する。このテーブルは、データが他のノードのライトで無効化された後も、新たなデータが同じアドレスに読み込まれるまでは関連するオーナシップ情報を残しておく。

Dynamic directory table に、要求されたデータのオーナシップ情報が無い場合、我々が fixed owner table と呼ぶもう一つのテーブルが必要になる。この状況はデータがローカル・メモリに取り込まれたことがないか、ローカル・メモリの同じアドレスに別のデータが取り込まれてオーナシップ情報が上書きされた場合に起こる。（我々はこの状況を cold cache miss と呼ぶ。これに対応して、ローカル・メモリにデータは残っていないが、dynamic directory table にオーナシップ情報が残っている状況を hot cache miss と呼ぶ。）全てのノードは fixed owner table を持ち、DSM 空間の予め割り当てられた範囲のデータのオーナシップを記憶している。各ノードに割り当てられる範囲の広さはそのノードのローカル・メモリの広さであり、全てのノードのローカル・メモリが同じ大きさならば、DSM 空間をノード数で割った大きさとなる。Cold cache miss が発生すると、データ要求メッセージは、まずマネージャに送られ、そこ

表 1: ノード#3 の Dynamic directory table

DSM Line	Phy. Line	Status	Owner	Fwd Link	Bwd Link
100	0	Invalid	5	—	—
101	1	exclusive	3	—	—
302	2	exclusive	3	—	—
—	3	—	—	—	—
—	999,999	Invalid	—	—	—
59	999,999	shared	10	2	4

DSM Line: DSM Line Number
 Phy. Line: Physical Line Number
 Status: Coherency Status
 Owner: Owner Node ID
 Fwd Link: Forward Link to Sharing Node
 Bwd Link: Backward Link to Sharing Node

から fixed owner table のエントリに記録されているオーナーへと転送される。Fixed owner table を用いたオーナーの探し出しは、マネージャを介するため、メッセージ数が増加する。しかし我々は cold cache miss 数は hot cache miss 数に比べて非常に少ないと考えている。このオーナーシップ管理アルゴリズムを semi-dynamic distributed manager algorithm と呼ぶことにする。

図 2 は、リクエスト、マネージャ、オーナーの 3 者の間のメッセージの流れを二種類示している。一方は hot cache miss の場合のもの、他方は cold cache miss の場合のものである。DSM アドレス#10 へのメモリ・アクセスは hot cache miss の場合のもので、リクエストの dynamic directory table に、オーナー・ノード ID#3 が記録されている。リクエストはデータ要求メッセージをオーナー・ノードに直接送る。オーナーは、DSM アドレス#10 を dynamic directory table を用いてローカル・メモリの物理アドレス#i に変換し、そこに置かれているデータ A をリクエストに送り返す。DSM アドレス#15 へのメモリ・アクセスは cold cache miss の場合のものである。リクエストは、まずデータ要求メッセージをマネージャに送る。マネージャは fixed owner table を参照して、オーナー・ノードの ID、#3 を発見し、メッセージをオーナーへと転送する。以降は hot cache miss 時と同様にして、データ B がオーナーによって供給される。

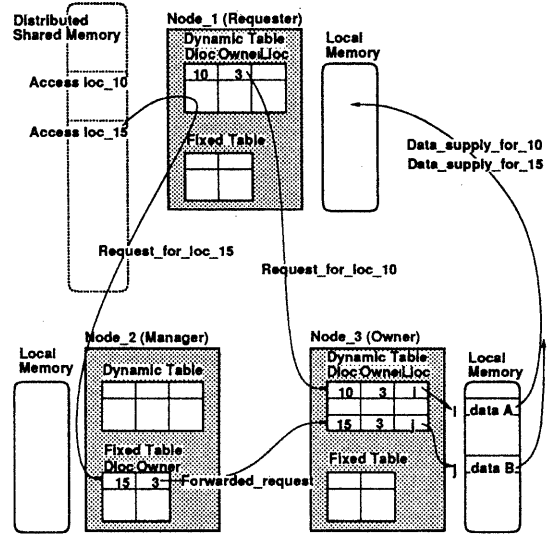


図 2: リクエスト、マネージャ、オーナーの 3 者の間のメッセージの流れ

4 Semi-dynamic distributed manager algorithm の性能

K. Li は文献 [4] において、fixed distributed manager algorithm と dynamic distributed manager algorithm のメッセージ・コストを、以下の式 1、および、式 2 によって示している。ここで、式中に現れる各変数は、次の意味を持っている。

- c_s : メッセージ送信コスト
- c_r : メッセージ受信コスト
- $c_{total}(i)$: プロセッサ i での page fault/cache miss 全体のコスト
- $c_{inv}(i)$: プロセッサ i での無効化の総コスト
- $f_r(i)$: プロセッサ i での read page fault/read cache miss の総数
- $f_w(i)$: プロセッサ i での write page fault/write cache miss の総数
- $r_p(i)$: プロセッサ i が受信したデータ要求の総数
- $r_{inv}(i)$: プロセッサ i が受信した無効化要求の総数
- $r_{loc}(i)$: プロセッサ i が受信したページ検索要求メッセージの総数

$s_{fwd}(i)$: プロセッサ i が転送した検索メッセージの総数

4.1 fixed distributed manager algorithm

$$\begin{aligned} C_{total}(i) = & c_{req1} + r_{loc}(i)(c_s + c_r) \\ & + r_p(i)(c_s + c_r) \\ & + c_{inv}(i) + r_{inv}(i)c_r \end{aligned} \quad (1)$$

但し,

$$c_{req1} = 2(f_r(i) + f_w(i))(c_s + c_r)$$

データを要求するのに必要なマネージャとオーナーの間で交わされるメッセージのコスト

$$r_{loc}(i) \approx \frac{1}{N} \sum_{j=1}^N (f_r(j) + f_w(j))$$

データを検索するのに必要なリクエストとマネージャの間で交わされるメッセージのコスト

4.2 dynamic distributed manager algorithm

$$\begin{aligned} C_{total}(i) = & c_{req2} + f_r(i) \times c_s + r_p(i) \times c_s \\ & + c_{inv}(i) + r_{inv} \times c_r \end{aligned} \quad (2)$$

但し,

$$c_{req2} = (f_r(i) + f_w(i))(c_s + c_r)$$

データを要求するのに必要なマネージャとオーナーの間で交わされるメッセージのコスト

$$\begin{aligned} \sum_{i=1}^N f_f(i) & \leq \sum_{i=1}^N s_{fwd}(i) \\ & \leq c(\log N) \sum_{i=1}^N (f_r(i) + f_w(i)) \end{aligned}$$

$f_f(i)$ は、ノード i において、データ要求メッセージを受信したが、自分のローカル・メモリにデータが存在しておらず、オーナーとして登録されているノードへ、そのメッセージを転送する総数

概略的に見て、無効化に必要なコストを除いたコストは、dynamic distributed manager algorithm の

方が fixed distributed manager algorithm の $\frac{1}{2}$ になることがわかる。semi-dynamic distributed manager algorithm のメッセージ・コストは、以下の式3によって示される。

4.3 semi-dynamic distributed manager algorithm

$$\begin{aligned} C_{total}(i) = & (1 - p_f(i)) \\ & \times (c_{req1} + (r_{loc}(i) + r_p(i))(c_s + c_r)) \\ & + p_f(i)(c_{req2} + f_f(i)c_s + r_p(i)c_s) \\ & + c_{inv}(i) + r_{inv}(i)c_r \end{aligned} \quad (3)$$

式3において、 $p_f(i)$ は、hot cache miss の確率である。この確率が1に近付くほど、式3の第2項の成分が大きくなり、その性能はdynamic distributed manager algorithm と変わらなくなる。例えば、キャッシュ・ライン長が128byteのシステム上で1Mbyteの子プロセスを実行することを想定する。単純に考えて、まず、8000回のcold cache missが発生する。しかし他の子プロセスと一箇所以上の共有領域を持ち、互いにこの領域に8000回の書き込み操作を行なったとすれば、 $p_f(i)$ の値は0.5となる。

以上のように、semi-dynamic distributed manager algorithm はシステム性能をdynamic distributed manager algorithm に近付けることができる。

5 キャッシュ・スペース割当プロトコル

COMAはライト・バック用の記憶は持たないので、ローカル・メモリ空間が満杯に近くなると、データ(DSM空間)のスラッシングが発生する。キャッシュ・スペースを再利用し、スラッシングによる性能の低下を抑えるためにOSによるdirectory tagの操作を伴うDSM空間中のデータの整理が必要となる。ユーザ・プロセスが解放したDSM空間を特定し、対応するfixed owner tableを持つマネージャにテーブルを更新させるメッセージを送る必要がある。以上の処理にはdirectory controllerと同様な

機能, すなわち DSM アドレスの物理アドレスへの変換, タグ情報の更新, そしてメッセージのやり取りが, OS が必要となる。

このハードウェア依存でタイミング・クリティカルな操作を OS に作り込むのを避けるために, 全 DSM 空間をローカル・メモリにマッピングして, かつ, 他ノードのデータをキャッシュする余裕を各ノードのローカル・メモリに残すようにハードウェアで自動的に制御するためのプロトコルについて考察する。DSM 空間をローカル・メモリにマッピングする場合, ノード間で共有されるデータがあるため, DSM 空間の容量はローカル・メモリの全容量よりも少なくなる。しかも共有データの量によって, ローカル・メモリにマッピングできる DSM 空間が変化するため, ハードウェアによる制御が困難になる。そこで, 各ノードが共有データ量の上限を定めることで, DSM 空間のローカル・メモリへのマッピングを directory controller が維持できるようになる。全 DSM 空間の大きさは, 以下の 3 つの定数と式 4 から求められる。

N : システム中のノード数

L_i : ノード i のローカル・メモリ・サイズ

S_i : ノード i の共有データ量の上限

$$\sum_{i=1}^N L_i(1 - S_i) + \max_{i=1}^N (L_i \times S_i) \quad (4)$$

この DSM 空間のサイズは各ローカル・メモリにキャッシュ・スペースを確保するために定義される。したがって OS は空きスペースを作り出すためにディレクトリ・タグを操作する必要がなくなる。各ノードはデータ・コヒーレンシとキャッシュ・スペースを, 以下の 6 つのローカル・レジスタを用いて式 5 を満たすように制御する。

$NODES$: システム中のノード数。上で N として与えられた。

$LSIZE$: ローカル・メモリのライン数。上で与えられた L_i に関係する。

$SDEGREE$: ノードの共有データの比率。上で S_i として与えられた。

$VCNT$: ローカル・メモリ中の有効なライン数

$SCNT$: ローカル・メモリ中の共有ライン数

$DSIZE$: DSM 空間中の総ライン数。式 4 で計算された。

$$VCNT \leq LSIZE$$

$$SCNT \leq LSIZE \times SDEGREE \quad (5)$$

$NODES, LSIZE$, と $SDEGREE$ は初期化の後には更新されないとする。すると $DSIZE$ も式 4 で求められた後は更新されない。 $VCNT$ と $SCNT$ は以下のデータ・コヒーレンシ・プロトコルに従って, 自動的に更新される。Directory controller は, このプロトコルを全てのリモート・アクセスに適用する。データ・コヒーレンシ・プロトコルは次の三状況で異なる。(1) 共有ラインが共有データ量の上限を越えている場合 (リモート・ノードからデータが要求された場合に発生する。), (2) 新しいデータを置くスペースがある場合, (3) 新しいデータを置くスペースがない場合。図 3 は, 各状況下で新しいデータをローカル・スペースに置く場合のデータ・コヒーレンシとキャッシュ・スペースを示している。

case 1: $SCNT \geq LSIZE \times SDEGREE$

ローカル・メモリ中の共有ライン数が上限を越えている場合である。Directory controller は共有ラインを上書きしてラインを再利用する候補として選び出す。候補のオーナーが自ノードならば, 他の共有ノードにオーナーシップを移す。Controller は自ノードを共有ノードのチェーンから外し, 新しいデータを上書きする。新しい共有ノードのチェーンに自ノードを入れる。

case 2: $(SCNT < LSIZE \times SDEGREE)$

$\cap (VCNT < LSIZE)$

ローカル・メモリに無効ラインがある場合。Directory controller は新しいデータを無効ラインに上書きし, 共有ノードのチェーンに自ノードを入れる。

case 3: $(SCNT < LSIZE \times SDEGREE)$

$\cap (VCNT = LSIZE)$

空きスペースがないが, 共有ライン数は上限を下回っている状態。Directory controller は専有ラインを選び出し, オーナシップと共に他ノードに送り出す。そのラインに新しいデータを上

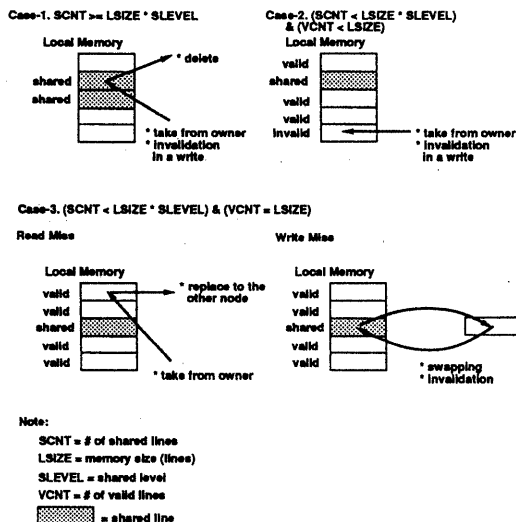


図 3: データ・コヒーレンシと共有データ量の上限の制御

書きし、共有ノードのチェーンに自ノードを入れる。

Directory controller は上で述べた6つのローカル・レジスタを調べてこれら三状況のいずれであるかを判断し、リモート・データを要求する前に空きスペースを確保する。

6 おわりに

本論文では、非階層ディレクトリを用いて flat COMA システムを実現するオーナーシップ管理機構、semi-dynamic distributed manager algorithm を提案し、簡単な評価を行なった。Semi-dynamic distributed manager algorithm は、ディレクトリの大きさが fixed distributed manager algorithm と同様にローカル・メモリの大きさによって決まり、かつ、システム性能を dynamic distributed manager algorithm に近付けることができる。また、COMA システムのデータ・スラッシングを防止するためのキャッシュ・スペース割当プロトコルについて考察した。

Semi-dynamic distributed manager algorithm, キャッシュ・スペース割当プロトコルについて、より詳細なモデル化、性能評価を行なうことが

今後の課題である。

参考文献

- [1] Gustavson, D. B. The Scalable Coherent Interface and Related Standards Projects. *IEEE Micro*, Vol. 12, No. 1, February 1992, pages 10-22.
- [2] Hagersten, E., Landin, A., and Haridi, S. DDM - A Cache-Only Memory Architecture. *IEEE Computer*, pages 44-54, September 1990.
- [3] Lenoski, D., Laudon, J., Gharachorloo, K., Gupta, A., Hennesy, J., Horowitz, M., and Lam, M. The Directory-Based Cache Coherence Protocol for the DASH Multiprocessor. In *Proceedings of the 17th Annual International Symposium on Computer Architecture*, pages 148-159, May 1990.
- [4] Li, K. and Hudak, P. Memory Coherence in Shared Virtual Memory Systems. *ACM Transactions on Computer Systems*, Vol. 7, No. 4, November 1989, pages 321-359.
- [5] Nitzberg, B. and Lo, V. Distributed Shared Memory: A Survey of Issues and Algorithms. *IEEE Computer*, Vol. 24, No. 8, August 1991, pages 52-60.
- [6] Stenström, P., Joe, T., and Gupta, A. Comparative Performance Evaluation of Cache-Coherent NUMA and COMA Architectures. In *Proceedings of the 19th Annual International Symposium on Computer Architecture*, pages 80-81, May, 1992.
- [7] X3T9.3 Task Group. Fibre Channel Physical and Signaling Interface (FC-PH). Rev. 2.2, January 1992.