

TR 907

スタック領域が不要な深さ優先順
コピー型ゴミ集め方式

中島 浩 (京都大学)

近山 隆

February, 1995

(C)Copyright 1995-2-17 ICOT, JAPAN ALL RIGHTS RESERVED

I C O T
Mita Kokusai Bldg. 21F
4-28 1-Chome
Minato-ku Tokyo 108 JAPAN

Tel(03)3456-3191~5

スタック領域が不要な深さ優先順コピー型ゴミ集め方式

中島 浩 (京都大学工学部) 近山 隆 (ICOT)

概要

代表的なゴミ集め方式のひとつである、コピー型のゴミ集めを改良した二つの方式を提案する。従来の方式が幅優先順にコピーを行なうのに対し、提案する方式はいずれも深さ優先順にコピーを行なう。この改良によって、ゴミ集め中およびゴミ集め後のメモリ・アクセスの局所性を大幅に改善できる。またデータ・オブジェクトだけを見て、その内容がポインタか否かを区別できる必要がないので、データ構造の設計に関する自由度が高い。

深さ優先順にコピーを行なうためには、未処理の要素を持つデータ構造を何らかの方法で記憶しておく必要がある。提案する方式の一つであるリンク法では、コピー前の領域に存在する未処理要素を持つデータ構造をリンクで結ぶことにより記憶する。またもう一つの方式である予約スタック法では、データ構造を指示するような未処理要素をコピー先領域の末端に配置されたスタックに記憶する。従って、いずれの方式においてもスタック領域を別途用意する必要はなく、従来の方式と同じ大きさのメモリ空間しか使用しない。

これら二つの方式と従来の方式の性能を評価し比較した所、ゴミ集めの対象となる領域が大きくなり、例えば実記憶容量を越えているような場合には、提案する方式ではページ・フォルトの回数が大幅に削減されることが明らかになった。

Depth-First Copying Garbage Collection without Extra Stack Space

Hiroshi Nakashima (Kyoto University) Takashi Chikayama (ICOT)

Abstract

In this paper, we propose two copying garbage collection methods. These two methods copy data structures in depth-first order, while conventional method copies in breadth-first order. This modification greatly improves memory access locality during both garbage collection and computation after. Restriction on data structure representation is also relaxed because the proposed methods don't require that a data object itself indicates whether it is a pointer.

For the depth-first copying, it is necessary to memorize information about non-leaf data structures so that their unprocessed elements are visited afterward. The first proposed method, the *link* method, links data structures which have unprocessed elements. Another method, the *reservation stack* method, pushes each unprocessed element being a pointer to another structure to a stack allocated at the end of the area to which data objects are copied. Therefore, neither of methods requires extra stack space.

We evaluated the performance of the proposed methods together with the conventional method. The result shows that the proposed methods have a great advantage over the conventional method when data objects are distributed in a large memory space, because the number of page faults are drastically reduced.

1 はじめに

自動メモリ領域管理機構は、ソフトウェア作成者からメモリ領域管理の負担を除き、より本質的な問題解決に専念させるために、殊に記号処理向きの言語の処理系では不可欠な機能となっている。

不要になったデータの占めるメモリ領域を自動的に再利用するゴミ集めには、さまざまな方式が提案され実際に用いられてきたが、代表的な方式の一つにコピー型のゴミ集め方式がある。これは、二つのメモリ領域を持ち、必要なデータだけを領域間でコピーすることによって、コピーされなかった不要データの占める領域を含め、コピー元の領域全体を再利用する方式である。この方式は不要データの占めるメモリ領域に一切アクセスしないため、必要なデータの割合が少ない場合に特に効率が良い。また、コピーの結果必要なデータがメモリ領域の一端に集まるので、連続した空き領域を提供でき、後のデータ割り付けが容易になり、ワーキングセットも小さくなるという長所を持つ。実装が比較的容易であることも大きな利点である。これらの長所は、メモリ領域が二つ必要なのでアドレス空間を2倍必要とするという欠点を十分に補うものであるため、最近の言語処理系ではコピー型の方式が多用されている。

さて、従来のコピー型ゴミ集め方式、即ち Fenichel と Yochelson によって最初に提唱された方式 [1] や、それを元に変形した諸方式 [2, 3] では、ネストしたデータ構造の領域間でのコピーが幅優先順に行なわれ、コピー後のデータも幅優先順に並ぶことになる。ところが、再帰的にデータ構造を処理するプログラムでは、データ構造を深さ優先順に扱うのが普通である。このため、データ構造はその作成時にはメモリ中に深さ優先順に並ぶように割り付けられ、後の利用でもこの順にアクセスされることが多い。このように深さ優先順に並んだデータ構造に対してゴミ集め時に幅優先順のコピーを行なうと、コピー時のワーキングセットが大きくなる。またコピー後の実行でも幅優先順に並んだデータを深さ優先順に用いることになるので、ワーキングセットが大きくなってしまう。

この問題を解決するために、本稿では深さ優先順にコピーを行なう二つの方式、リンク法と予約スタック法を提案する。深さ優先順にコピーを行なうためには、未処理の要素を持つデータ構造を何らかの方法で記憶しておく必要があるが、リンク法ではこのようなデータ構造をリンクで結ぶことにより記憶する。またもう一つの方式である予約スタック法では、データ構造を指示するような未処理要素をコピー先領域の末端に配置されたスタックに記憶する。従って、いずれの方式においてもスタック領域を別途用意する必要はなく、メモリ量やゴミ集めに要する手間のオーダは従来の方式と等しい。また、ゴミ集めのための特殊なデータ・タグを設けたり、各メモリ語に余分なビットを付加する必要もない。更に、提案する方式では常にデータ構造を指すポインタを用いながらコピーを行なうので、データ・オブジェクトだけを見て、その内容がポインタか否かを区別できる必要がない。

以下本稿では、まず2章において従来のコピー型ゴミ集め方式とその長所短所について簡単に述べる。次に3章では、提案する二つの幅優先順ゴミ集め方式について説明し、4章ではこれらの方式と従来の方式の性能比較を行なう。続いて5章で関連する研究について述べた後、最後に6章で本稿の結論を述べる。

2 従来のコピー型ゴミ集め方式

提案する方式の説明の前提を与えるため、本章ではまず従来のコピー型ゴミ集め方式について概説する。

2.1 ゴミ集めの手順

コピー型ゴミ集め方式の基本的な手順は以下の通りである。

- (1) 同じ大きさのメモリ領域をふたつ用意する。
- (2) 通常のメモリ割り付けにはその一方だけを用いる。
- (3) 使用中のメモリ領域（以下この領域を旧領域と呼ぶ）がいっぱいになったら、必要なデータだけをもう一方の領域（新領域）にコピーする。不要なデータはコピーしないので、コピー後の領域には空きができる。
- (4) 以降、コピーした先の領域をメモリ割り付けに用いる。

本稿で提案する方式も、このレベルでの手順は従来のコピー型ゴミ集め方式とまったく異ならない。違いはステップ (3) のコピーを行なう方式の詳細がけである。

従来の方式では、所謂 Cheney のアルゴリズム [4] に基づき、概略以下の手順にしたがってコピーを行なう（詳細は付録 A.2 参照）。なおデータ領域やデータ構造はアドレスが小さい方から大きい方へ向かって連続して伸びることを前提とし、その最も小さいアドレスの語を先頭、最も大きいアドレスの語を末尾と呼ぶ。またポインタ p が指す語を $*p$ と表記する。

[B1] ルートのコピー

新領域の末尾を指すポインタ new_tail を新領域のベースとする。次に全てのルート（実行に直接必要とわかっているポインタ群）について、各々のアドレスをポインタ new に代入した上で [B3] の手続きを実行する。その後 new に新領域のベースを代入する。

[B2] 新領域のスキャン

new と new_tail が等しくなるまで、 new をインクリメントしながら [B3] の手続きを繰り返し実行する。

[B3] コピー手続き

[B3-1] ポインタの識別

$*new$ がポインタ語でなければ手続きを終了する。

[B3-2] コピー済の判定

S を $*new$ が指示するデータ構造とする。 S の先頭語が新領域を指示するポインタ語であれば、 S は先頭語が指す領域にコピー済である。そこで、 $*new$ がコピー先を指すように変更し、手続きを終了する。

[B3-3] データ構造のコピー

S の全体を new_tail から始まる領域にコピーし、 $*new$ がコピー先を指すように修正する。

[B3-4] コピー済の通知

S の先頭語にコピー先を指すポインタ語を格納して、データ構造がコピーされたことを示す。その後で new_tail をデータ構造の語数だけ増す。

さてこの手順では、新領域の new と new_tail の間を、要素が処理されていないデータ構造を保持する FIFO キューとしており、ネストしたデータ構造は幅優先でコピーされる。また [B2] で new を順次増やししながら新領域をスキャンし、[B3-1] で各々の語がポインタ語か否かを識別していることにも注意が必要である。

この他、データ構造とそれを指すポインタ語を調べることにより、以下の二つが明らかになるというを前提を用いているが、これらは方式に関わらず必要なものであり、かつ妥当なものである。

(p1) データ構造の大きさ。

(p2) データ構造がコピー済であることを示す新領域へのポインタ語の位置。

なお (p2) の語には、本来の値として新領域へのポインタ語と解釈される値が入っていない。また前述の手順と以下の説明では、データ構造やポインタ語に関わらず先頭語を用いているが、一般にはどの語であっても良い。更に、複数語のデータ構造へのポインタ語はその先頭を指すことによってデータ構造の全体を指示すること、即ち：

(p3) データ構造の一部分のみを指すようなポインタは存在しない。

ことを前提としており、この制限は提案する方式についても同様である。この制限の緩和については紙数の都合で省略するが、[5] で論じているように

- 論理型言語の処理系で用いられる変数参照ポインタのような「透明な」ポインタの導入
- コピー済のデータ構造の先頭語を識別する特殊タグの導入

によって対処することができる。

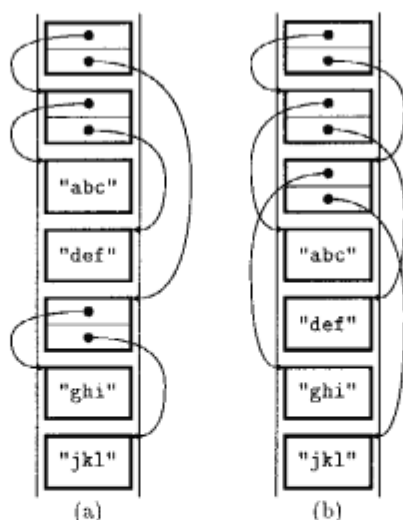


図 1: 従来方式での 2 進木のコピー

2.2 従来方式の性質

従来のコピー型ゴミ集め方式は、以下の長所を持っている。

- (+1) 不要なデータが占めるメモリ領域に一切アクセスしないので、必要なデータの割合が少ない場合に効率が良い。
- (+2) ゴミ集めが必要なデータをメモリ領域の一端に集めるので；
 - ゴミ集め後に空き領域が連続領域になり、データ割り付けが容易になる。
 - コピー後のデータをアクセスする際に、データを移動しないゴミ集め方式に比べて、ワーキングセットが小さくなる。
- (+3) 算法が簡潔で実装が比較的容易である。

一方、メモリ領域を二つ用いるために 2 倍のアドレス空間が必要とするというコピー型に共通の欠点があるほか、従来の方式に本質的な問題点として以下の二つが挙げられる。

- (-1) ネストしたデータ構造については、幅優先順にコピーを行なうが、これは通常の再帰呼び出しなどによるプログラムのアクセス順である深さ優先順と一致していない。例えば Prolog において；

```
tree([L|R]):- left(L), right(R).
left(["abc" | "def"]).
right(["def" | "ghi"]).
```

によって四つの文字列を葉に持つ 2 進木構造；

```
[[["abc" | "def"] | ["ghi" | "jkl"]]
```

を生成すると、図 1(a) のように配列する¹。ところがこれを幅優先順にコピーすると、新領域では図 1(b) のような配列となる。このため、元のデータの並び順を保存する方式、例えばスライディング・コンパクション方式[6]に比べると；

- 深さ優先順に作成したデータ構造を、ゴミ集めでコピーするとき
- ゴミ集めでコピーしたデータ構造を、深さ優先順にアクセスするとき

¹多くの Prolog 処理系では "abc" などの文字列を文字コードのリストとして表現するが、ここでは文字列特有のデータ構造があるものとして説明する。

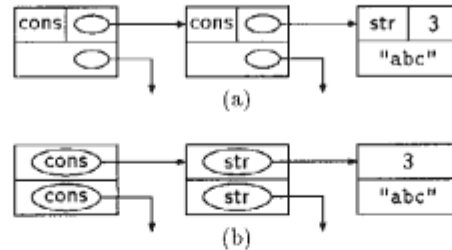


図 2: オブジェクト・タグとポインタ・タグ

に参照アドレスの局所性が低い。

- (-2) 新領域をスキャンしながらコピーすべきデータ構造を見つけるため、ポインタとそうでないデータを識別するようなデータ表現が必要である。即ち、ある語がポインタ語か否かがデータ構造自体を調べるだけで明らかになる必要がある。この条件はデータ表現形式の設計に制約を加えるものであり、ゴミ集めの際に知る必要があるデータ構造の性質を、それを指すポインタのみが記述するような設計ができない。

例えば前述の 2 進木の例では、“abc”などの文字列をその長さで表す文字コードの並びで表現しようとする、それを誤ってポインタと解釈しないようにするための措置が必要となる。その一つとして図 2(a) のように、cons セルや文字列の先頭にデータ構造の型を示すタグを付け、cons セルでは 2 つのポインタを処理し、文字列では長さ分のデータをスキップする方法がある。この方法でのタグは、付加された語が属するオブジェクトの性質を示すためオブジェクト・タグと呼ばれる。しかしオブジェクト・タグは型識別のためのオブジェクトの参照、即ちデータ型を知るためにポインタが指示するメモリ語を読む必要がある。従って、図 2(b) のようにタグをポインタに付け、ポインタが指すオブジェクトの性質を示すようにするポインタ・タグに比べて一般に効率が悪い。

3 深さ優先順のコピー型ゴミ集め方式

従来の方式が持つ二つの問題点を解決するためには、深さ優先順にコピーすることと、新領域の単純なスキャンを不要とすることの、二つの改善が必要である。この二つは、ネストしたデータ構造のコピーのために新領域を FIFO キューとして用いることに密接に関連している。従って別途 LIFO スタックを設け、ネストしたデータ構造を深さ優先順にコピーすることにすれば、二つの改善を同時に実現できる。

例えば、複数の要素を持ち、かつ各々が他のデータ構造へのポインタでありうるようなデータ構造（以下非終端構造と呼ぶ） S の要素 e が、他の非終端構造 T を指すポインタである時、以下の処理文脈をスタックに退避した後で T の処理を行えば良い。

- (1) S を指示するポインタ語。
- (2) e の次の要素を知るための情報（ S の先頭からのオフセットなど）。

スタックは T の処理が完了した時点でポップされ、 e の次の要素から S の処理が再開される。その際、 S の大きさやコピー先はスタックに退避した S へのポインタ語から求めることができる。またスタックが空であることは、一つのルートに関する処理が完了したことを意味する。

この方法（単純スタック法と呼ぶ）に基づくコピーの手順は付録 A.3 に示すが、本質的な部分は以下の通りである。

[D1] 文脈を保持する変数

非終端構造 S の末尾要素ではない要素 e を処理しているとき、処理の文脈は S を指すポインタ語 $pword$ と、 S の先頭から e へのオフセット $offset$ が保持する。

[D2] その他の変数

新領域の末尾を new_tail とする。またコピーしようとする旧領域の語の値を $word$ 、そのコピー先アドレスを new とする。

[D3] データ型の識別

$word$ のデータ型に基づき、以下の処理を行なう。

[D3-1] 非ポインタのコピー

$word$ がポインタ語でなければ $*new$ に $word$ を格納し、[D4] に飛ぶ。

[D3-2] コピー済の判定

$word$ が指示するデータ構造を T とする。 T の先頭語が新領域を指示するポインタ語であれば、データ構造はコピー済またはコピー中である。そこで $*new$ に $word$ と同じデータ型でコピー先を指すポインタを格納し、[D4] に飛ぶ。

[D3-3] ヘッダのコピー

$*new$ に $word$ と同じデータ型で new_tail を値とするポインタを格納し、 T にヘッダがあればそれを new_tail から始まる領域にコピーする。次に $word$ にヘッダの次の語を代入し、 T の先頭語にコピー先へのポインタを格納する。また

$$\begin{aligned} new &\leftarrow new_tail + \langle T \text{ のヘッダの語数} \rangle; \\ new_tail &\leftarrow new_tail + \langle T \text{ の語数} \rangle; \end{aligned}$$

として、 T のコピー先を確保する。

[D3-4] ヘッダのみの構造

T のヘッダを除く語数が 0 であれば、ポインタを持たない構造であるので [D4] に飛ぶ。

[D3-5] 1 要素の構造

ヘッダを除く語数が 1 であれば T の文脈を保存する必要がないので [D3] に戻る。

[D3-6] 非終端構造

ヘッダを除く語数が 2 以上であれば $pword$ と $offset + 1$ をスタックにプッシュした後

$$pword \leftarrow word; \quad offset \leftarrow \langle \text{ヘッダの語数} \rangle;$$

として T の処理文脈を設定する。[D3] に戻る。

[D4] 次要素の処理

[D4-1] スタックの空判定

スタックが空であれば一つのルートに関する処理が終了。

[D4-2] 次要素の設定

$offset$ に 1 を加え、 $pword$ と $offset$ より次の要素の値とそのコピー先アドレスを求め、 $word$ と new に代入する。

[D4-3] 末尾要素判定

末尾要素であればスタックをポップして文脈である $pword$ と $offset$ を復元する。

[D4-4] 次要素の処理

次の要素を処理するために [D3] に戻る。

なお説明を簡単にするために、上記の手順ではデータ構造の表現形式を以下のように仮定している。

(p4) データ構造は先頭に連続した非ポインタ語からなるヘッダを持つことができ、その大きさはデータ構造とそれを指すポインタ語を調べることにより明らかになる。

(p5) データ構造のヘッダを除いた各語はポインタ語である可能性があり、ポインタ語か否かはその語のみを調べれば明らかになる。

但しこれらの前提は単純スタック法にとって本質的ではなく、データ構造 S の語 w がポインタ語か否かを、 S へのポインタと S の先頭からの m のオフセットから知ることができれば充分である。

さてこの手順では非終端構造 S の要素を処理中に、未コピーの非終端構造 T へのポインタを発見すると、[D3-6]において S の処理文脈をスタックにプッシュし、処理文脈を T のものに更新して先頭要素の処理

を開始する。一方要素が非ポインタ語、コピー済またはコピー中のデータ構造へのポインタ、あるいは非ポインタ語のみからなるデータ構造であれば、[D4]により S の次の要素の処理に移る。但しそれが末尾要素であれば S の処理に戻る必要はないので、その時点でスタックをポップして S の「親」の文脈を復元する。この一種の tail recursion 操作は要素数が1のデータ構造の処理にも適用され、文脈の退避/復元は行なわない。

単純スタック法の問題点は、スタックの深さが最悪の場合には旧領域と同じサイズになり、コピー型の欠点である大きなアドレス空間を更に拡大してしまうことにある。そこで、必要なメモリ領域は従来の方式と同一に保ったまま、深さ優先順にコピーする方式を二つ提案する。

3.1 リンク法

データ構造の表現形式に関する前述の仮定 (p4) と (p5) が成り立つ場合、非終端構造 S の要素 e の処理文脈として、以下を用いることができる。

- (1) e の旧領域アドレス (*old_link*)
- (2) e の新領域アドレス (*new_link*)
- (3) 未処理要素数

リンク法では、これらの文脈を図3に示すようにコピー中の非終端構造をリンクで結ぶ形で保存し、スタックを用いずに深さ優先順のコピーを行なう。即ち e が非終端構造 T へのポインタである時、単純スタック法の処理ステップ [D3-6] を以下の手順に置き換える（詳細は付録 A.4 参照）。

[D3-6'] 非終端構造

old_link を T の末尾要素に保存し、末尾要素自身は T の新領域の末尾に退避する。また e の値自身の参照が完了していることを利用し、*new_link* を $*old_link$ （旧領域の e ）に保存する。次に T のヘッダの次の語に関する旧領域/新領域のアドレスを、それぞれ *old_link* と *new_link* に代入する。

一方、未処理要素数については、等価な情報として S の末尾要素が「親」の非終端構造へのリンクであることを利用する。即ち、このリンクと他の要素の値とを区別できるようにし、これをセンチネルとして用いて一つの非終端構造の処理が完了したことを知る。この区別のために特殊なタグや余分なビットを用いなくても済むように、本来は旧領域のアドレスである *old_link* を、以下の式を用いて新領域アドレス *old_link'* に変換し、ポインタであることを示す適当なタグを付加して退避する。

$$old_link' = old_link - old_base + new_base$$

但し *old_base* と *new_base* はそれぞれ旧領域/新領域のベースである。この変換によって、最終要素の値は新領域へのポインタに見えるが、前提 (p5) と 2.1 で述べた前提 (p3) により、このようなポインタは末尾要素ではない未処理の要素の値として出現することはない。従って単純スタック法の処理ステップ [D4-2 ~ 3] を以下の手順に置き換えればよい（詳細は付録 A.4 参照）。

[D4-2'] 次要素の設定

old_link と *new_link* にそれぞれ1を加え、次の要素 $*old_link$ を *word* に代入する。また *new* にコピー先アドレスである *new_link* を代入する。

[D4-3'] 末尾要素判定

word が新領域へのポインタ語であれば、末尾要素であるので文脈を復元する。即ち

$old_link \leftarrow (word \text{ を旧領域アドレスに変換した値 }) ;$
 $word \leftarrow *new_link ; \quad new_link \leftarrow *old_link ;$

なお空のスタック、即ち非終端構造の連鎖の末端を示す空ポインタは、旧領域中に出現しない任意のアドレスで表現することができ、初期化の際にその値を *old_link* に格納すれば良い。但しその場合には *new_link* の退避と復元、及び末尾要素の判定において、空ポインタに関する特別な操作が必要となる。一方、新旧各々の領域に1語ずつ、前述の変換式で相互に変換できるようなアドレスを確保し、この使用しないアドレスを空ポインタに用いれば、これらの特別な操作は不要となる。

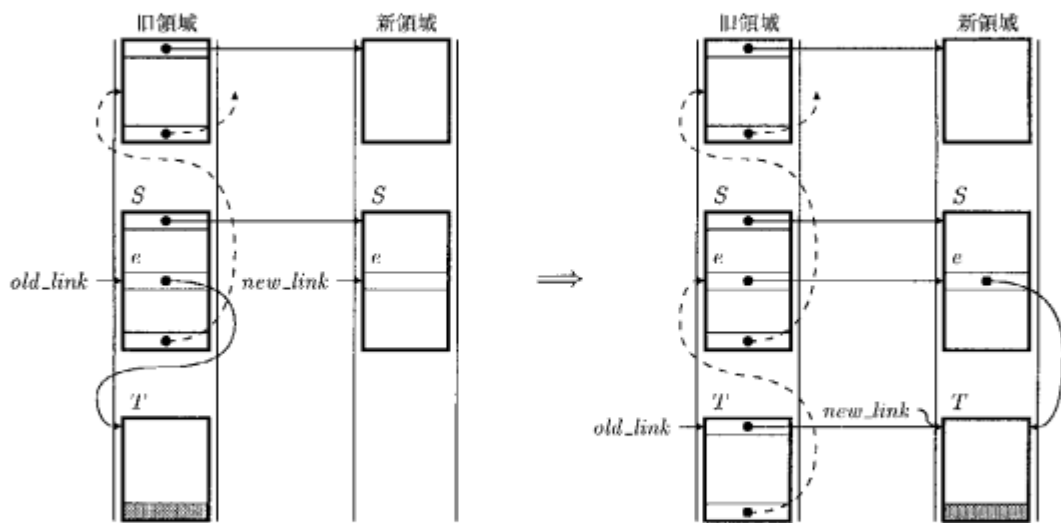


図 3: 非終端構造のリンク

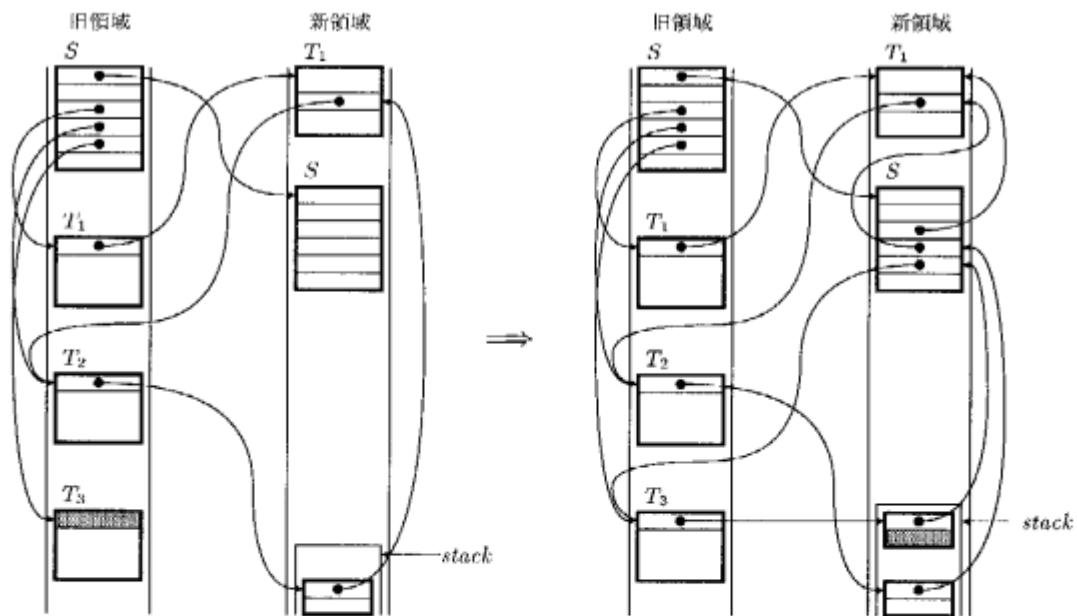


図 4: 予約スタック

3.2 予約スタック法

前述の単純スタック法では、非終端構造 S の要素として非終端構造 T が出現した時に、 S の処理を中断して T の処理を行ない、スタックには S の処理に復帰するための情報を退避した。一方、提案する予約スタック法では、 S の処理を中断せずに T を処理するための情報をスタックに退避し、 S の処理完了後に T を処理する。この『予約スタック』は新領域のベースと反対側の末端に設け、コピーによって新領域の末尾 *new_tail* が伸びる方向と逆の方向に伸びるようにし、余分なスタック領域を使用しないで済ませることができる。

予約スタックに退避する情報が、非終端構造一つ当たり 2 語以下であり、かつ退避されるのが一度だけであれば、新領域の末尾と予約スタックが衝突しないことを保証できる。即ち、 T の情報をスタックに退避する時点では、 T のコピー先は確保されていないが、 T がその後にコピーされることは確定している。また T は非終端構造であるので、明らかに 2 語以上の領域を占めている。従って、既にコピーされたデータ構造の語数とスタックの深さの和は、最終的にコピーされるべき全てのデータ構造の語数を越えることはなく、新領域の大きさを越えることも当然ない。

具体的には以下の 2 語をスタックに退避すれば良い。

- (1) T を指すポインタ t のコピー先アドレス
- (2) T の先頭語の値

t のコピー先アドレスは、予約スタックへの退避の時点では T のコピー先は確定していないために保存する必要がある。即ち、 T を発見した時点では t を単に新領域にコピーし、後に T を処理する際に T のコピー先を指すように修正する。なお、 T 自身のアドレスは t が保持しているので、予約スタックから間接的に知ることができ、退避する必要はない。また後述するように、 T が実際にコピーされる前に T を指す他のポインタが見つかった場合には、そのような全てのポインタのコピー先が連鎖を形成し、予約スタックが連鎖の先頭を、また連鎖の末端が T を指すようにする。

一方、 T に関する情報が一度がけしか予約スタックに退避されないようにするためには、 T が「予約済」であることを何らかの方法で通知する必要がある。これを特殊なタグや余分なビットを用いずに行なうために、 T の先頭語を退避先のアドレスに変更し、かつその元の値を予約スタックに退避する。予約スタックは新領域に存在するので、非終端構造の先頭語の値によって、未コピー、コピー済、コピー予約済のいずれであるかを、以下のように判定することができる (図 4)。

- 新領域へのポインタでなければ未コピー (T_0)
- 新領域へのポインタであるが、予約スタックへのポインタでなければコピー済 (T_1)
- 予約スタックへのポインタであれば、コピー予約済 (T_2)

また、非終端構造 T がポインタ t の処理によってコピー予約され、その後 T を指す別のポインタ t' が見つかった場合、 T が処理される時に t と t' の双方が修正されるようにしなければならない。そこで、図 4 に示すように：

$$stack \rightarrow \langle t' \text{ のコピー先} \rangle \rightarrow \langle t \text{ のコピー先} \rangle \rightarrow T$$

の連鎖を形成し、 T の処理時に連鎖中の全てのポインタが T のコピー先を指すように修正できるようにする。なお連鎖中のポインタは新領域またはルート領域にあるので、旧領域にある T を指すポインタが出現すれば連鎖の末端であると判定できる。

以上をまとめると、ある語 *word* を新領域アドレス *new* にコピーする際に、*word* が 2 語以上のデータ構造 T を指すポインタである時の手順は以下のようになる。

[R1] S がコピー済

*new に *word* と同じデータ型で T のコピー先を指すポインタを格納する。

[R2] S がコピー予約済

T の先頭語の値、即ち T の情報が退避されている予約スタック中のアドレスを *sentry* とし、

$$*new \leftarrow *sentry ; *sentry \leftarrow new$$

として、*word* のコピー先をポインタの連鎖の先頭に挿入する。

[R3] T が未コピー

new と T の先頭語の値を予約スタックにプッシュし、 $*new$ に $word$ を、また T の先頭語に予約スタック・トップへのポインタを、それぞれ格納する。

なお非ポインタ語と 1 語のデータ構造を指すポインタについては、単純スタック法と基本的に同じ処理を行えば良い。

また、あるデータ構造の全ての要素のコピーが完了したならば、以下の処理を行なう。

[R4] スタックの空判定

スタックが空であれば一つのルートに関する処理が終了。

[R5] スタック・ポップ

予約スタックをポップし、コピー予約されたデータ構造を指すべきポインタの連鎖の先頭アドレスと、データ構造の先頭語の値を得る。

[R6] ポインタ連鎖の修正

旧領域へのポインタが出現するまで連鎖をたどり、連鎖中の全てのポインタがデータ構造のコピー先 (new_tail) を指すようにする。

[R7] データ構造のコピー開始

連鎖の末端にある旧領域のポインタが指すコピー予約したデータ構造を T とする。 T の先頭語にコピー先へのポインタ (new_tail) を格納し、

$$\begin{aligned} new &\leftarrow new_tail; \\ new_tail &\leftarrow new_tail + (T \text{ の語数}); \end{aligned}$$

として T のコピー先を確保した後、予約スタックから得た T の先頭語の元の値から T の各要素のコピーを開始する。

以上をまとめた手順の詳細を付録 A.5 に示す。なおこの手順では前述の仮定 (p4) と (p5) を用いているが、これらの仮定は予約スタック法にとって本質的ではない。また、説明を簡単にするために非終端構造の処理を先頭要素から正順に行なっているが、コピーの順序を左優先にするために末尾から逆順に行なうこともできる。

4 性能評価

本章では従来の方式と提案した二つの方式の性能を、メモリのアクセス回数、小～中規模のメモリ領域を対象とした時の実行時間、及び大規模なメモリ領域を対象とした時のページ・フォルトの回数の、二つの指標に基づいて比較する。

4.1 メモリ・アクセス回数

従来の方式では、データ構造を旧領域から新領域へコピーした後で、新領域をスキャンしてその要素を処理する。従ってデータ構造中のポインタとなりうる要素の各々について 3 回のメモリ・アクセス、即ちコピーのための旧領域での読出と新領域への書込、及びスキャンによる新領域での読出が行なわれる。一方、提案した二つの方式はいずれも新領域のスキャンを行なわないので、要素ごとのメモリアクセス回数は旧領域での読出と新領域への書込の 2 回で済む。従って少なくとも要素数が充分大きければ、提案した方式の方が従来の方式よりも高速であると言える。

実際にはコピー済通知などのために、要素数によらずデータ構造ごとに一定回数のメモリ・アクセスを必要とし、この回数は方式によって異なる。例えば、ただ一つのポインタから指され、かつポインタ語を 2 語以上含むようなデータ構造に関して、各々の方式における定数回のメモリ・アクセスは以下ようになる。

従来の方式 = 2 回

- コピー済通知のための先頭語の書込 (1 回)
- データ構造を指すポインタの修正 (1 回)

リンク法 = 7回

- コピー済通知のための先頭語の書込 (1回)
- 末尾語の退避のための読出/書込 (各1回)
- *old_link* の退避/復元のための書込/読出 (各1回)
- *new_link* の退避/復元のための書込/読出 (各1回)

予約スタック法 = 8回

- コピー済通知のための先頭語の書込 (1回)
- コピー予約済通知のための先頭語の書込 (1回)
- スタックの書込/読出 (各2回)
- データ構造を指すポインタ修正のための読出/書込 (各1回)

従って、全ての語がポインタ語でありうるような n 語のデータ構造の場合、各方式でのメモリ・アクセス回数は：

$$\begin{aligned}\text{従来方式} &= 3n + 2 \\ \text{リンク法} &= 2n + 7 \\ \text{予約スタック法} &= 2n + 8\end{aligned}$$

となり、従来方式に比べて $n > 5$ でリンク法が、また $n > 6$ でスタック法が、それぞれメモリ・アクセス回数が少なくなる。

4.2 実行時間

ゴミ集めに要する時間を決定する要因として、前節で述べたメモリ・アクセス回数の他に、処理の複雑さとメモリ・アクセスの局所性などが挙げられる。提案した二方式は従来方式に比べ、処理の複雑さの面では劣り、メモリ・アクセスの局所性の面では優るものと考えられる。これらを総合的に評価するため、各方式に基づくゴミ集めのプログラムを作成し、実行時間を測定した。

測定に用いたデータは $n = 2^k$ のバランスした n 進木である。木の各ノードである n 要素のデータ構造は $n+1$ 語からなり、先頭語に要素数 n が、引き続く n 語に各要素がそれぞれ格納される。この先頭語のコピーのためのメモリ・アクセス回数の増加は従来法で3、リンク法とスタック法では2である。但し Lisp や Prolog などの言語では、 $n = 2$ の cons セルを2語で実装する方式が良く用いられるので、それに準拠した表現とした。即ち cons セルに限り、データ構造を指すポインタを他のものと区別し、要素数を省略している。また木の葉には非ポインタ語が直接格納されている。1語は32ビットであり、下位2ビットがデータ型を示すタグ、上位30ビットがアドレスなどに用いられる。

ゴミ集めの方法は基本的に2章と3章で述べたものに基づいているが、リンク法と予約スタック法に関しては、tail recursion 最適化の技法を用いて不要なメモリ・アクセスを削除している。その結果評価に用いた n 進木ではメモリ・アクセス回数が若干減少し、 $n > 2$ の場合のノードあたりの平均アクセス回数は：

$$\begin{aligned}\text{リンク法} &\approx 2n + 7 + 2(n-1)/n \\ \text{予約スタック法} &\approx 2n + 3 + 7(n-1)/n\end{aligned}$$

と近似できる。また cons セルの場合には近似的に、リンク法が9.5、予約スタック法が8.5となる。更に、予

約スタック法は左優先にコピーするために要素の処理を逆順に行なっている。

また、プログラムにはゴミ集めを行なう部分の他に n 進木を操作する部分があり、木を左優先/深さ優先に探索しながら、木のコピーを生成する。例えば2進木の場合、以下の Prolog プログラムと等価な操作を行なう。

```
traverse([L1|R1],[L2|R2]):-!,
    traverse(L1,L2),traverse(R1,R2).
traverse(T,T).
```

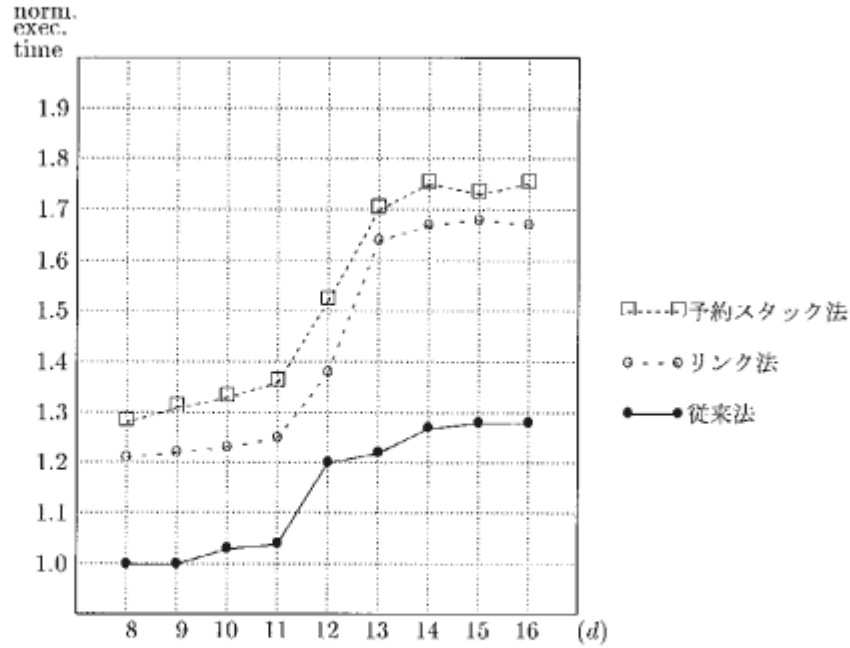


図 5: 2 進木の場合の実行時間

実行時間の測定はこの操作も含めて行なった。

プログラムは C 言語によって作成し、主記憶容量 16 MB、キャッシュ容量 64 KB の SparcStation1 で実行した。またコンパイラは gcc 2.4.5 版を用い、コンパイル・オプションは“-O2 -fomit-frame-pointer”である。

図 5 は、深さ d が 8 から 16 の 2 進木の操作とゴミ集めに要した時間を木の語数で割り、それを更に深さ 8 の時の従来方式の実行時間を 1 として正規化したものである。なお深さを 17 以上にすると特に従来方式でページ・フォルトが頻発し、有意な測定ができなかった。提案した手法は従来方式に比べ実行時間が長くなっているが、 d が小さい時にはリンク法で約 20%、予約スタック法で約 30% の性能低下に留まっている。リンク法については、従来法に比べてメモリ・アクセス回数がノードあたり 1.5 回 (18.8%) 増加しているのが性能低下の主な原因である。一方、予約スタック法ではメモリ・アクセス回数増はノードあたり 0.5 回 (6.3%) であるので、実行時間差の多くの部分は処理の複雑さによるものである。またいずれの方式でも深さが 12 になると急に実行時間が増加するが、これはデータがキャッシュからあふれることが原因である。従来法の方が増加の割合が少ない理由は、新領域のスキャンの際にキャッシュのプリフェッチ効果が現れ、キャッシュ・ミスが頻度が少なくなることであると考えられる。

一方、図 6(a) は n を変化した場合の性能であり、 d はページ・フォルトが頻発して測定不能になる直前まで大きくした。なお、 $n = 2, d = 16$ の時の従来法の実行時間を 1 として正規化してある。また同図 (b) は 1 語当たりのメモリ・アクセス回数の実測値を、 $n = 2, d = 16$ における従来法の値を 1 として正規化したものである。図から明らかなように、 n が増加するに従ってリンク法や予約スタック法の性能が相対的に良くなり、従来法と逆転する。この理由は前述のメモリ・アクセス回数の差が主なものであるが、リンク法や予約スタック法でもキャッシュのプリフェッチ効果が、木の葉の部分のコピーの際に現れることにも関係していると考えられる。

なおバランスした n 進木では、後述するように従来法のメモリ・アクセス局所性が提案した方式に比べて非常に悪い。そこで、アクセス局所性がほぼ同一となる線形リストについても性能評価を行なった。評価に用いたリストは、car 部を非ポインタ語とし、長さは 2^n ($n = 8 \sim 16$) とした。1 語あたりの実行時間はどの方法についても長さに依存せずほぼ一定であり、従来法を 1 とするとリンク法は約 1.1、また予約スタック法は約 1.5 であった。リンク法と予約スタック法の差が大きいのは、リンク法が $\text{car} \rightarrow \text{cdr}$ と処理するため tail recursion 最適化の効果が現れるのに対し、予約スタック法では $\text{cdr} \rightarrow \text{car}$ と処理するために効果がないことによる。実際、従来法と比べたノードあたりのアクセス回数増は、リンク法が 1 であるのに対し、予約スタック法は 4 となってしまう。また予約スタック法での要素処理順序を正順とすると、tail recursion 最適化の効

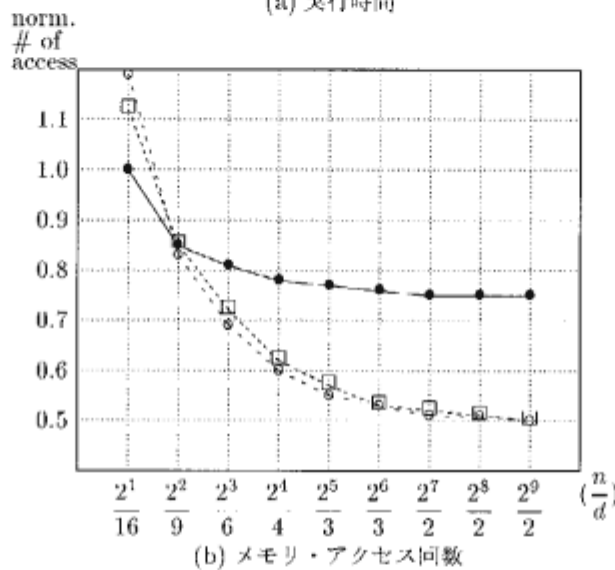
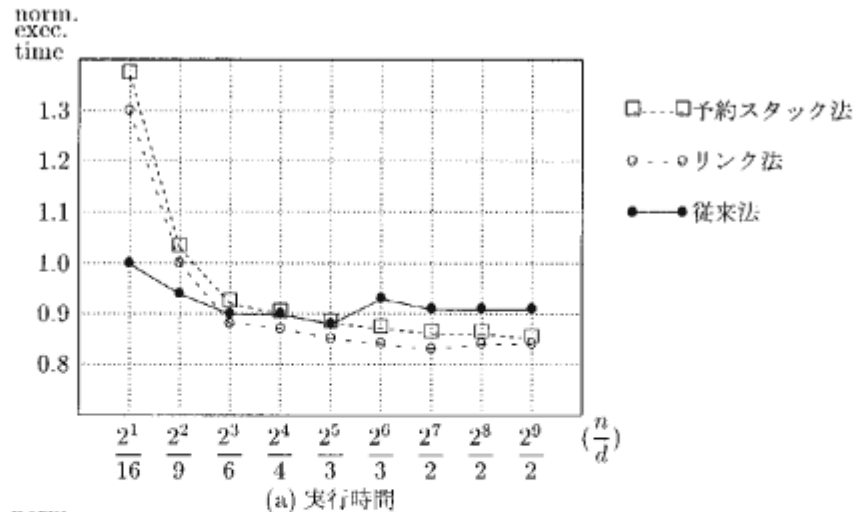


図 6: n 進木の場合の実行時間とメモリ・アクセス回数

果によってノードあたりのアクセス回数が従来法よりも 3 回減少し、従来法を 1 とした実行時間も約 0.87 となった。

4.3 ページ・フォルトの回数

従来法の問題点であるメモリ・アクセス局所性の悪さは、ページ・フォルトが増加する形で顕在化するものと考えられる。即ち、幅優先で肝領域をアクセスするために、連続してアクセスする二つのデータ構造間のアドレス差が大きく、これらが同じメモリ・ページに存在する確率が低くなることが予想される。

そこで、各方式について深さが 20 の 2 進木 (約 8 MB) を対象として、ページ・フォルトがどのように発生するかを評価した。但しこの処理を実際の計算機上で行なうと、ページ・フォルトの頻発のために有意な測定ができないので、アクセスするメモリ・アドレスのトレースを生成してページ・フォルトの発生状況をシミュレートした。なお、メモリ・ページの大きさは 4 KB とし、ゴミ集め開始時にはアクセス対象の全てのページがスワップ・アウトされているものと仮定した。ゴミ集めでは新旧各々の領域について 8 MB = 2048 ページをアクセスするため、4096 回のページ・フォルトが必然的に発生する。そこで、これを基準として更に余分なページ・フォルトがどれだけ発生するかを、物理ページ数を変化させながら測定した。

まず表 1 は、余分なページ・フォルトが全く発生しないようにするために必要な物理ページ数を示したものである。表から明らかなように、リンク法ではアクセスするページの約 1/4、また予約スタック法では約 1/2

表 1: 余分なページ・フォルトを起こさないための物理ページ数

従来法	リンク法	予約スタック法
3585	1026	2050

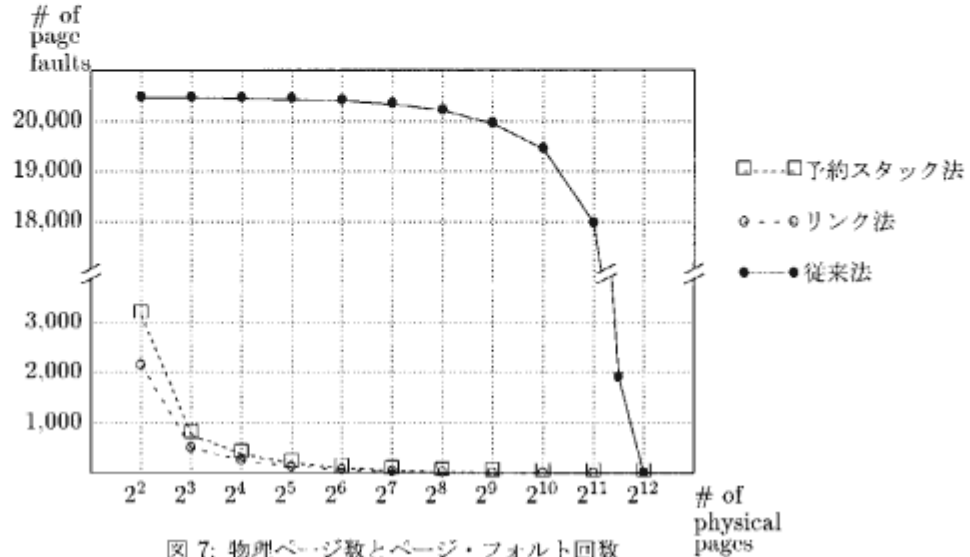


図 7: 物理ページ数とページ・フォルト回数

の物理ページがあれば、余分なページフォルトは全く発生しない。これに対して従来法では約 7/8 の物理ページが必要であり、アクセス局所性が悪いことが明らかになった。また、ゴミ集めの結果を深さ優先/左優先で探索する場合、リンク法や予約スタック法では 1 ページあれば余分なページ・フォルトは全く発生しないが、従来法では実に 2047 ページが必要であることも明らかになった。

次に、物理ページ数を表 1 に示した値よりも少なくした時に、余分なページ・フォルトが何回発生するかを測定した。結果は図 7 に示す通りであり、リンク法や予約スタック法では物理ページ数の減少に対するページ・フォルトの増加は緩やかである。これに対し従来法では急激に増加し、アクセスするページ数の 1/2 である 2048 ページでも、実に約 18,000 回ものページ・フォルトが発生する。この原因は、ある一定の深さ以下では同じ深さにあるノードが全て別のページに存在し、アクセスするノードが変わるたびにページ・フォルトとなってしまうことにある。

以上のように従来法ではアクセスする領域が大きくなるとページ・フォルトが極めて高い頻度で発生するため、実用的な時間でゴミ集めができないことが明らかになった。また、新旧領域の合計を物理メモリ容量以下にしたとしても、ページ・アクセス局所性の悪さは TLB のヒット率低下につながり、やはり大幅な性能低下を招くことは明らかである。

5 関連研究

深さ優先順にコピーを行なう比較的に容易な方法として、3 章で述べたような単純なスタックを用いるもの、あるいは Deutsch-Schorr-Waite の逆転ポインタ法 [7, 8] を用いるものが考えられ、例えば [9] などでは幅優先順に比べて良好なアクセス局所性が得られることが示されている。しかし前者はスタックのためにメモリ領域を必要とし、後者は各語にゴミ集め専用の特別なマークビットを必要とする。

またアクセス局所性を改善する別の方法として、Wilson らは 2 レベルの幅優先順コピー法を提案している [10]。この方法では新領域のページ p をスキャン中に、 $p' \neq p$ であるようなページ p' にデータ構造をコピーすると、 p のスキャンを中断し、コピーしたデータ構造をルートとした p' 内のスキャンを行なう。このスキャンの中断は再帰的に発生するため、元のページに復帰するための情報を保持する 1 エントリが 2 語のスタック

が必要である。このスタックの深さは新領域のページ数と同じだけあれば充分ではあるが、余分なメモリ領域を必要とすることには変わりはない。また4.3節で述べたような大規模なデータ構造を深さ優先順に生成／参照すると、やはり大量のページ・フォルトが発生するものと予想される。

一方、コピー型ゴミ集めの問題点として、旧領域におけるデータ構造の生成順序が新領域では保存されないというものがある。この問題点の解決については提案した方式が直接貢献することはないが、例えば[11]に示されている Prolog の heap 領域のセグメント生成順序を保存する方法に組み込んで、スタック領域やマークビットを除去することはできる。またこれに関連した研究として、小出らによる任意のデータ構造の生成順序を保存するゴミ集め方式がある[12]。しかしこの方法では、ゴミではないデータをマークするためにスタックなどが必要であり、またマークビットを不要にするためにはポインタが指示するデータ構造が2語以上であるという前提が必要である。

6 おわりに

本稿ではコピー型のゴミ集めを改良し、深さ優先順にコピーを行なう二つの方式を提案した。いずれの方式についても、不要なデータにアクセスしないことや、必要なデータがメモリ領域の一端に集まるなど、従来方式の利点をそのまま受け継いでいる。また、ゴミ集めに要する手間のオーダや必要とするメモリ量も従来方式と等しく、特殊なデータ・タグや各メモリ語に余分なビットを必要としないことは、関連する他の研究には見られない大きな特質である。

一方、従来方式が幅優先順にコピーをするために生じる参照アドレス局所性の低さは、深さ優先順にコピーすることにより解消しており、ページ・フォルト回数の劇的な減少と言う評価結果によって提案した方式の有効性を示すことができた。また、新領域のスキャンを必要としないため、ポインタ・タグを使用できるなど、従来方式に比べてデータ構造の設計に関する自由度が大幅に増したことも、提案した方式の大きな特質である。

参考文献

- [1] Fenichel, R. and Yochelson, Y.: A Lisp Garbage Collector for Virtual Memory Computer Systems, *Comm. ACM*, Vol. 12, No. 11, pp. 419-429 (1969).
- [2] Baker, Jr., H. G.: List Processing in Real Time on a Serial Computer, *Comm. ACM*, Vol. 21, No. 4, pp. 280-294 (1978).
- [3] Unger, D.: *The Design and Evaluation of a High Performance Smalltalk System*, ACM Distinguished Dissertations, The MIT Press (1987).
- [4] Cheney, C. J.: A Nonrecursive List Compacting Algorithm, *Comm. ACM*, Vol. 13, No. 11, pp. 677-678 (1970).
- [5] 中島浩, 近山隆: データ構造の一部を指すポインタを許容するコピー型ゴミ集め方式, Technical memorandum, ICOT (1994). (to be published.).
- [6] Morris, F. L.: A Time and Space Efficient Garbage Collection Algorithm, *Comm. ACM*, Vol. 21, No. 8, pp. 662-665 (1978).
- [7] Cohen, J.: Garbage Collection of Linked Data Structure, *Computing Surveys*, Vol. 13, No. 3, pp. 341-367 (1981).
- [8] Schorr, H. and Waite, W. M.: An Efficient Machine-Independent Procedure for Garbage Collection in Various List Structures, *Comm. ACM*, Vol. 10, No. 8, pp. 501-506 (1967).
- [9] Stamos, J. W.: Static Grouping of Small Objects to Enhance Performance of a Paged Virtual Memory, *ACM Trans. on Programming Languages and Systems*, Vol. 2, No. 2, pp. 155-180 (1984).
- [10] Wilson, P. R., Lam, M. S. and Moher, T. G.: Effective "Static-graph" Reorganization to Improve Locality in Garbage-Collected Systems, *Proc. ACM SIGPLAN'91*, pp. 177-191 (1991).

- [11] Bekkers, Y., Ridoux, O. and Ungaro, L.: Dynamic Memory Management for Sequential Logic Programming Languages, *Proc. Intl. Workshop on Memory Management 92*, LNCS 637, Springer-Verlag, pp. 82–102 (1992).
- [12] 小出洋, 野下浩平: 生成順序を保存するコピー方式ガベージコレクションについて, 情報処理学会論文誌, Vol. 34, No. 11, pp. 2395–2400 (1993).

付録

A.1 データ型, 変数, 手続き/関数の定義

```

define mem_word = 〈メモリ語を表すデータ型〉;
define mem_addr = 〈メモリ語のアドレスを表すデータ型〉;
define mem_tag = 〈メモリ語のタグを表すデータ型〉;

old_base : mem_addr ;      {旧領域の先頭アドレス}
new_base : mem_addr ;      {新領域の先頭アドレス}
new_size : integer ;        {新領域の語数}
set_of_roots : set of mem_addr ; {ルートのアドレスの集合}
ptag : mem_tag ;           {ポインタ型のタグ}

function load(a : mem_addr) : mem_word ; {a の語の値を返す}
procedure store(a : mem_addr, w : mem_word) ; {a に w を書き込む}
function tag(w : mem_word) : mem_tag ; {w のタグを返す}
function address(w : mem_word) : mem_addr ; {w が指す語のアドレスを返す}
function make_word(l : mem_tag, a : mem_addr) : mem_word ;
    {タグが l, アドレスが a のポインタ語を返す}
function is_pointer(w : mem_word) : boolean ; {w がポインタ語であれば真}
function copied(w : mem_word) : boolean ; {w が指すデータ構造がコピー済であれば真}
function object_size(w : mem_word) : integer ; {w が指すデータ構造の語数を返す}
procedure copy_words(ao : mem_addr, n : integer, an : mem_addr) ; {ao から n 語を an にコピー}

```

A.2 従来の方式の手順

```

procedure breadth_gc ;
    new_tail : mem_addr ; { (新領域の末尾)+1 }
    new : mem_addr ; { 処理対象語の新領域アドレス }
    word : mem_word ; { 処理対象語の値 }
    old : mem_addr ; { ポインタ語が指す語の旧領域アドレス }
    size : integer ; { ポインタ語が指すデータ構造の語数 }

    procedure copy_data ; { データのコピー }
    begin
        word ← load(new) ;
        if is_pointer(word) then begin { ポインタ語の処理 }
            old ← address(word) ;
            if copied(word) then { オブジェクトはコピー済 }
                store(new, make_word(tag(word), address(load(old)))) ;
            else begin
                store(new, make_word(tag(word), new_tail)) ; { ポインタ語がコピー先を指すようにする }
                size ← object_size(word) ;
                copy_words(old, size, new_tail) ; { オブジェクトを新領域にコピー }
            end
        end

```

```

        store(old, make_word(ptag, new.tail)) ;      {先頭語にコピー済を通知}
        new_tail ← new_tail + size ;                {コピー先の確保}
    end ;
end ;
end ;
begin
    new_tail ← new_base ;
    for ∀new ∈ set_of_roots do                      {全てのルートが指すデータをコピー}
        copy_data ;
        new ← new_base ;
        while new < new_tail do begin              {コピーされたデータが指すデータをコピー}
            copy_data ;
            new ← new + 1 ;
        end ;
    end ;
end ;

```

A.3 単純スタック法の手順

```

procedure depth_gc ;
    { breadth_gc と同じ変数宣言 }
    pword :      mem_word ; {コピー中非終端構造へのポインタ語}
    offset :      integer ; {処理対象要素のオフセット}
    old_head :    mem_addr ; {コピー中非終端構造の先頭アドレス (旧領域)}
    new_head :    mem_addr ; {コピー中非終端構造の先頭アドレス (新領域)}
    n :           integer ; {未処理語数}
    hsize :       integer ; {ヘッダの語数}
    stack_empty : boolean ; {スタックが空ならば true}

    function copy_header(w : mem_word, a : mem_addr) ;
        {w が指すデータ構造のヘッダを a にコピーしヘッダの語数を返す}

    procedure initialize_stack ;                      {スタックの初期化}
    begin
        stack ← stack_base ;
        pword ← NULL ;
        offset ← 0 ;
    end ;

    procedure push_stack ;                          {スタックのプッシュ}
    begin
        store(stack, pword) ;                      {非終端構造へのポインタ語を退避}
        store(stack + 1, offset + 1) ;              {次要素のオフセットを退避}
        stack ← stack + 2 ;
        pword ← word ;                             {非終端構造へのポインタ語と先頭アドレスを設定}
        offset ← hsize ;
        old_head ← old ;
        new_head ← new - offset ;
        n ← size - offset ;                         {未処理要素数を設定}
    end ;

    procedure next_element ;                        {次要素の処理}
    begin
        if stack = stack_base then                {スタックは空}
            stack_empty ← true ;

```

```

else begin
    offset ← offset + 1 ;
    word ← load(old_head + offset) ;
    new ← new_head + offset ;
    n ← n - 1 ;
    if n = 1 then begin
        stack ← stack - 2 ;
        pword ← load(stack) ;
        offset ← load(stack + 1) ;
        if pword ≠ NULL then begin
            old_head ← address(pword) ;
            new_head ← address(load(old_head)) ;
            n ← object_size(pword) - offset ;
        end ;
    end ;
end ;
end ;
end ;
begin
    new_tail ← new_base ;
    initialize_stack ;
    for ∀new ∈ set_of_roots do begin
        word ← load(new) ;
        repeat begin
            stack_empty ← false ;
            if is_pointer(word) then begin
                old ← address(word) ;
                if copied(word) then begin
                    store(new, make_word(tag(word), address(load(old)))) ;
                    next_element ;
                end ;
            else begin
                size ← object_size(word) ;
                store(new, make_word(tag(word), new_tail)) ;
                hsize ← copy_header(word, new_tail) ;
                new ← new_tail + hsize ;
                word ← load(old + hsize) ;
                store(old, make_word(ptag, new_tail)) ;
                new_tail ← new_tail + size ;
                if size = hsize then
                    next_element ;
                else if size - hsize > 1 then
                    push_stack ;
                end ;
            end ;
        end ;
        store(new, word) ;
        next_element ;
    end ;
end ;

```

{コピー中の非終端構造あり}
 {次要素の設定}
 {未処理要素数を減らす}
 {末尾要素}
 {スタックをポップ}
 {非終端構造へのポインタ語を復元}
 {要素のオフセットを復元}
 {非終端構造の先頭アドレスを復元}
 {未処理要素数を復元}

{スタックを初期化}
 {全てのルート进行处理}
 {ルートを処理}
 {ポインタ語の処理}
 {オブジェクトはコピー済}
 {ポインタがコピー先を指すようにする}
 {ヘッダをコピーし先頭要素のオフセットを得る}
 {先頭要素のコピー先を設定}
 {先頭要素をロード}
 {先頭語にコピー済を通知}
 {コピー先の確保}
 {非ポインタ語のみからなるデータ構造}
 {非終端構造}
 {上記以外であれば1要素のデータ構造}

{非ポインタ語の処理}

```

    until stack_empty ;                                {スタックが空になるまで繰り返し}
end ;
end ;
end ;

```

A.4 リンク法の手順

```

procedure link_gc ;
    { breadth_gc と同じ変数宣言 }
    old_base :      mem_addr ; { 旧領域のベース }
    old_link :      mem_addr ; { 旧領域へのリンク }
    new_link :      mem_addr ; { 新領域へのリンク }
    hsize :         integer ;   { ヘッダの語数 }
    stack_empty :   boolean ;   { スタックが空ならば true }

    function copy_header(w : mem_word, a : mem_addr) ; { depth_gc と同じ }
    function old_to_new(a : mem_addr) ; old_to_new  $\leftarrow a - \text{old\_base} + \text{new\_base}$  ;
    function last_word(w : mem_word) ; { w が新領域へのポインタ語ならば真 }
    procedure initialize_stack ;                                { スタックの初期化 }
    begin
        old_link  $\leftarrow \text{NULL}$  ;                                { リンクを空ポインタに初期化 }
    end ;
    procedure push_stack ;                                    { スタックのプッシュ }
    begin
        store(new_tail - 1, load(old + size - 1)) ;        { 末尾要素の退避 }
        store(old_link, new_link) ;                          { リンクの退避 }
        store(old + size - 1, make_word(ptag, old_to_new(old_link))) ;
        old_link  $\leftarrow \text{old} + \text{hsize}$  ;                      { リンクを先頭要素に設定 }
        new_link  $\leftarrow \text{new}$  ;
    end ;
    procedure next_element ;
    begin
        if old_link = NULL then                                { スタックは空 }
            stack_empty  $\leftarrow \text{true}$  ;
        else begin                                              { コピー中の構造体あり }
            old_link  $\leftarrow \text{old\_link} + 1$  ;                    { リンクを次要素に設定 }
            new_link  $\leftarrow \text{new\_link} + 1$  ;
            new  $\leftarrow \text{new\_link}$  ;                                { 次要素の設定 }
            word  $\leftarrow \text{load}(\text{old\_link})$  ;
            if last_word(word) then begin                    { 末尾要素 }
                old_link  $\leftarrow \text{new\_to\_old}(\text{address}(\text{word}))$  ; { 旧領域リンクの復元 }
                word  $\leftarrow \text{load}(\text{new\_link})$  ;                { 末尾要素の復元 }
                new_link  $\leftarrow \text{load}(\text{old\_link})$  ;            { 新領域リンクの復元 }
            end ;
        end ;
    end ;
    begin
        { depth_gc と同じ手続き }
    end ;

```

A.5 予約スタック法の手順

```

procedure stack_gc ;
  { breadth_gc と同じ変数宣言 }
  old_e :      mem_addr ; { 処理対象要素の旧領域アドレス }
  new_e :      mem_addr ; { 処理対象要素の旧領域アドレス }
  hsize :      integer ;   { ヘッダの語数 }
  n :          integer ;   { 未処理語数 }
  stack_empty : boolean ;   { スタックが空ならば true }

  function copy_header(w : mem_word, a : mem_addr) ; { depth_gc と同じ }
  function old_area(a : mem_addr) ; { a が旧領域アドレスであれば真 }
  function reserved(w : mem_word) ; { w が指すデータ構造がコピー予約済であれば真 }
  procedure pop_stack ;
  begin
    new ← address(load(stack)) ;
    repeat begin                                     { 連鎖中の全てのポインタがコピー先を指すようにする }
      word ← load(new) ;
      store(new, make_word(tag(word), new_tail)) ;
      new ← address(word) ;
    end ;
    until old_area(new) ;                             { 旧領域アドレスが見つかるまで繰り返し }
    old ← new ;
    store(old, load(stack + 1)) ;                       { 先頭語を復元 }
    size ← object_size(word) ;
    hsize ← copy_header(word, new_tail) ;             { ヘッダをコピーし先頭要素のオフセットを得る }
    new ← new_tail + hsize ; new_e ← new ;             { 先頭要素のコピー先を設定 }
    old_e ← old + hsize ;
    word ← load(old_e) ;                                 { 先頭要素をロード }
    store(old, make_word(ptag, new_tail)) ;           { 先頭語にコピー済を通知 }
    new_tail ← new_tail + size ;                         { コピー先の確保 }
    n ← size - hsize ;                                   { 未処理要素数を設定 }
    stack ← stack + 2 ;                                  { スタックをポップ }
  end ;
  procedure next_element ;
  begin
    n ← n - 1 ;                                         { 未処理要素数を減らす }
    if n > 0 then begin                                { 未処理要素あり }
      old_e ← old_e + 1 ;                               { 次の要素を処理 }
      new_e ← new_e + 1 ; new ← new_e ;
      word ← load(old_e) ;
    end ;
    else if stack = new_base + new_size                { スタックは空 }
      stack_empty ← true ;
    else
      pop_stack ;                                       { スタックをポップ }
    end ;
  begin
    new_tail ← new_base ;
    stack ← new_base + new_size ;                     { スタックを初期化 }
    for  $\forall new \in set\_of\_roots$  do begin

```

```

word ← load(new) ;                                {ルートを処理}
n ← 1 ;
repeat begin
  stack_empty ← false ;
  if is_pointer(word) then begin                    {ポインタ語の処理}
    old ← address(word) ;
    if copied(word) then begin                      {オブジェクトはコピー済}
      store(new, make_word(tag(word), address(load(old)))) ;
      next_element ;
    end ;
    else if reserved(word) then begin               {オブジェクトはコピー予約済}
      old ← address(load(old)) ;                     {スタックのアドレス}
      store(new, make_word(tag(word), address(load(old)))) ;
      store(old, make_word(ptag, new)) ;              {ポインタを連鎖の先頭に挿入}
      next_element ;
    end ;
    else if object_size(word) = 1 then begin        {1語のオブジェクト}
      store(new, make_word(tag(word), new_tail)) ; {ポインタがコピー先を指すようにする}
      word ← load(old) ;                             {ポインタの先を続けて処理}
      store(old, make_word(ptag, new_tail)) ;
      new ← new_tail ;
      new_tail ← new_tail + 1 ;
    end ;
    else begin                                       {2語以上のオブジェクト}
      store(new, word) ;                             {ポインタ語を単にコピー}
      stack ← stack - 2 ;
      store(stack, make_word(ptag, new)) ;            {ポインタのアドレスを回避}
      store(stack + 1, load(old)) ;                  {先頭語を回避}
      store(old, make_word(ptag, stack)) ;            {先頭語にコピー予約済を通知}
      next_element ;
    end ;
  end ;
  else begin                                       {非ポインタ語の処理}
    store(new, word) ;
    next_element ;
  end ;
end ;
until stack_empty ;                               {スタックが空になるまで繰り返し}
end ;
end ;

```