

(19) 日本国特許庁(JP)

(12) 公開特許公報(A)

(11) 特許出願公開番号

特開2016-115288

(P2016-115288A)

(43) 公開日 平成28年6月23日(2016.6.23)

(51) Int.Cl.	F I	テーマコード (参考)
G 0 6 F 17/30 (2006.01)	G 0 6 F 17/30 4 1 4 Z	
	G 0 6 F 17/30 1 7 0 A	

審査請求 未請求 請求項の数 12 O L (全 24 頁)

(21) 出願番号	特願2014-255623 (P2014-255623)	(71) 出願人	390009531
(22) 出願日	平成26年12月17日 (2014.12.17)		インターナショナル・ビジネス・マシーンズ・コーポレーション INTERNATIONAL BUSINESS MACHINES CORPORATION アメリカ合衆国10504 ニューヨーク州 アーモンク ニュー オーチャードロード New Orchard Road, Armonk, New York 10504, United States of America
		(74) 代理人	100108501 弁理士 上野 剛史

最終頁に続く

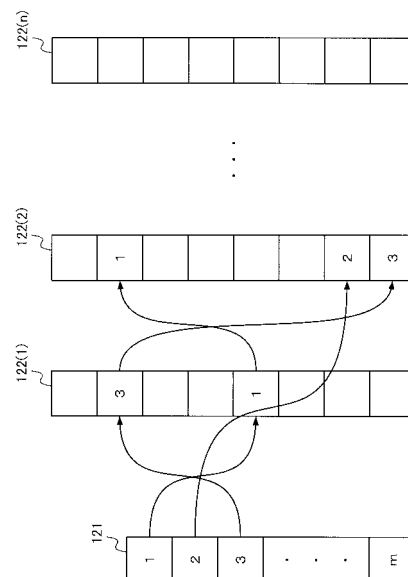
(54) 【発明の名称】 データを集計するシステム、方法およびプログラム

(57) 【要約】

【課題】集計処理において、必要なメモリ領域（記憶容量）を削減し、かつ、対象データや属性の更新処理に要する手間を削減する。

【解決手段】データを集計するシステムであって、対象データごとに所定の順序で並び、対象データの識別情報から対象データが有する属性のリストを指定するD t o Kインデックス122と、複数の対象データが有する属性のリストである単語リスト121と、を記憶するインデックス記憶部と、属性ごとにこの属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計する集計処理を行う集計処理部と、を備える。そして、単語リスト121の各属性に関して、対象データごとの第1リストの要素を順に辿るリンクが設定されており、集計処理部は、属性ごとに設定されたリンクに基づき、属性を有する対象データを調べる。

【選択図】図2



【特許請求の範囲】**【請求項 1】**

データを集計するシステムであって、

集計処理の対象である対象データの識別情報から当該対象データが有する属性のリストである第 1 リストを指定するインデックスであって、当該対象データごとの当該第 1 リストは所定の順序で並び、当該第 1 リストの各要素が当該属性の情報を含む、当該インデックスと、複数の当該対象データが有する属性のリストである第 2 リストと、を記憶するインデックス記憶部と、

前記属性ごとに当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計する集計処理を行う集計処理部と、を備え、

前記第 2 リストの各属性に関して、前記対象データごとの前記第 1 リストの要素を順に辿るリンクが設定されており、

前記集計処理部は、前記属性ごとに設定された前記リンクに基づき、当該属性を有する前記対象データを調べる、システム。

【請求項 2】

前記リンクは、前記対象データごとの前記第 1 リストの各要素が、当該要素に含まれるのと同じ前記属性の情報を含む最も近い後続の第 1 リストの要素を特定する情報を含むことにより、設定される、請求項 1 に記載のシステム。

【請求項 3】

前記第 2 リストの各属性は、複数の前記対象データにおける個々の対象データが当該属性を有する頻度の降順にソートされている、請求項 1 に記載のシステム。

【請求項 4】

前記対象データごとの前記第 1 リストの各要素は、当該要素が含む情報により特定される前記属性に対して付与された値をさらに含み、

前記第 1 リストの各要素は、前記値に基づく順序でソートされている、請求項 1 に記載のシステム。

【請求項 5】

前記属性に対して付与された値は、当該属性と当該属性を有する前記対象データとの関係に基づく重み値であり、

前記第 1 リストの各要素は、前記値の降順にソートされている、請求項 4 に記載のシステム。

【請求項 6】

前記対象データごとの前記第 1 リストは、最も古い第 1 リストを最後尾として作成された順に並び、新規に作成された第 1 リストを追加する場合は、第 1 リストの並びの先頭に配置され、かつ、新規に作成された当該第 1 リストの各要素に含まれる属性に関するリンクが更新される、請求項 1 に記載のシステム。

【請求項 7】

前記対象データごとの前記第 1 リストは、最も古い第 1 リストを最後尾として作成された順に並び、最も古い第 1 リストを削除する場合は、最後尾の第 1 リストが削除され、かつ、削除された当該第 1 リストの各要素に含まれる属性に関するリンクが更新される、請求項 1 に記載のシステム。

【請求項 8】

データを集計するシステムであって、

集計処理の対象である対象データごとに当該対象データが有する属性を登録したリストを含むインデックスを記憶するインデックス記憶部と、

前記属性ごとに当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第 1 手法、および、前記対象データごとに当該対象データが有する前記属性を調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第 2 手法のどちらで集計処理を行うかを判定する判定部と、

前記判定部による判定に応じて前記第 1 手法および前記第 2 手法の一方により集計処理

10

20

30

40

50

を行う集計処理部と、を備え、

前記インデックスは、相異なる前記対象データが有する同一の前記属性どうしを関連付けており、

前記集計処理部は、前記第１手法により集計処理を行う場合に、前記インデックスにおける同一の前記属性どうしの関連に基づき、当該属性を有する前記対象データを調べる、システム。

【請求項９】

データを集計する方法であって、

集計処理の対象である対象データの識別情報から当該対象データが有する属性のリストである第１リストを指定するインデックスであって、当該対象データごとの当該第１リストは所定の順序で並び、当該第１リストの各要素が当該属性の情報を含む、当該インデックスと、複数の当該対象データが有する属性のリストである第２リストと、前記第２リストの各属性に関して、前記対象データごとの前記第１リストの要素を順に辿るリンク構造と、を記憶するインデックス記憶部と、

10

前記属性の集計処理を行う集計処理部と、を備えるシステムにおいて、

前記集計処理部が、

前記属性ごとに設定された前記リンクに基づき、当該属性を有する前記対象データの数を調べるステップと、

前記属性を、調べた対象データの数が多い上位の所定数の属性を集計し、集計結果を出力するステップと、

20

を含む、方法。

【請求項１０】

データを集計する方法であって、

集計処理の対象である対象データごとに当該対象データが有する属性を登録したリストを含むインデックスを記憶するインデックス記憶部と、

前記属性の集計処理を行う集計処理部と、

前記集計処理部により行われる集計処理の手法を判定する判定部と、を備えるシステムにおいて、

前記判定部が、前記属性ごとに当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第１手法、および、前記対象データごとに当該対象データが有する前記属性を調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第２手法のどちらで集計処理を行うかを判定するステップと、

30

前記判定部が前記第１手法により集計処理を行うと判定した場合に、前記集計処理部が、相異なる前記対象データが有する同一の前記属性どうしを関連付けてなる前記インデックスにおける同一の前記属性どうしの関連に基づき、当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計するステップと、

前記判定部が前記第２手法により集計処理を行うと判定した場合に、前記集計処理部が、前記インデックスに基づき、各対象データが有する前記属性を調べ、当該対象データとの関係が所定の基準を満たす属性を集計するステップと、

を含む、方法。

40

【請求項１１】

コンピュータを、

集計処理の対象である対象データの識別情報から当該対象データが有する属性のリストである第１リストを指定するインデックスであって、当該対象データごとの当該第１リストは所定の順序で並び、当該第１リストの各要素が当該属性の情報を含む、当該インデックスと、複数の当該対象データが有する属性のリストである第２リストと、を記憶するインデックス記憶手段と、

前記属性ごとに当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計する集計処理を行う集計処理手段として機能させ、

前記第２リストの各属性に関して、前記対象データごとの前記第１リストの要素を順に

50

辿るリンクが設定されており、

前記集計処理手段の機能として、前記属性ごとに設定された前記リンクに基づき、当該属性を有する前記対象データを調べる、プログラム。

【請求項 12】

コンピュータを、

集計処理の対象である対象データごとに当該対象データが有する属性を登録したリストを含み、相異なる前記対象データが有する同一の前記属性どうしを関連付けてなるインデックスを記憶するインデックス記憶手段と、

前記属性ごとに当該属性を有する前記対象データを調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第1手法、および、前記対象データごとに当該対象データが有する前記属性を調べ、当該対象データとの関係が所定の基準を満たす属性を集計する第2手法のどちらで集計処理を行うかを判定する判定手段と、

前記判定手段による判定に応じて前記第1手法および前記第2手法の一方により集計処理を行う集計処理手段として機能させ、

前記集計処理手段の機能として、前記第1手法により集計処理を行う場合に、前記インデックスにおける同一の前記属性どうしの関連に基づき、当該属性を有する前記対象データを調べる、プログラム。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、データを集計するシステム、方法およびプログラムに関し、特に所定の属性を有する対象データの集合中から所定の属性を集計する技術に関する。

【背景技術】

【0002】

複数種類の属性を持つ大量のデータに対する集計操作の一つに、処理対象として指定されたデータ（以下、対象データ）の集合中に出現する属性を、ある規則に基づいて各属性に付与された値である点数（重み、スコア）の高いものから順にk種類求める集計操作がある。この集計操作は、Top-k等とも呼ばれる。例えば、対象データが文書（テキスト）で、属性がその文書に含まれる単語である場合を考える。また、属性である各単語に付与される点数を全文書中の出現数とする。このとき、特定の条件を満たす複数の文書だけに出現する単語のうち、出現頻度上位k個の単語を求める操作は、Top-kの集計操作の一例である。

【0003】

この種の従来技術として、テキストの識別情報から当該テキストに含まれるキーワードのリストを指定する第1インデックス（DOC__TO__KEYインデックス）と、キーワードの識別情報から当該キーワードを含むテキストのリストを指定する第2インデックス（KEY__TO__DOCインデックス）とを用いて、テキストを検索する技術がある（特許文献1参照）。ここで、テキストは上記の「対象データ」の一例であり、キーワードは上記の「属性」の一例である。この従来技術では、検索条件の入力を受け付けると、まず、第1インデックスによる検索時間および第2インデックスによる検索時間を見積もる。そして、より高速であると判断されたインデックスを用いて検索する。

【先行技術文献】

【特許文献】

【0004】

【特許文献1】特開2007-156739号公報

【発明の概要】

【発明が解決しようとする課題】

【0005】

KEY__TO__DOCインデックス（以下、KtoDインデックスと略記する）を用いた処理は、例えば、次のような手順による処理となる。すなわち、まず、出現頻度の高い

10

20

30

40

50

順に属性を順次選択する。そして、選択された属性を含む対象データのリストが、検索条件を満たすか否かを判断する。そして、検索条件を満たす対象データの数が多い順にk個の属性を選択する。これにより、選択された属性が集計結果となる。しかし、この処理は、集計される対象となる属性の種類が多い場合、多大な処理時間を要する。

【0006】

DOC__TO__KEYインデックス(以下、D to Kインデックスと略記する)を用いた処理は、例えば、次のような手順による処理となる。すなわち、まず、検索条件を満たす対象データを選択する。そして、選択された対象データの識別情報に対応する属性のリストを取得する。そして、リストに示された属性を含む対象データの数を属性毎に集計する。しかし、この処理は、対象データの数が多い場合、多大な処理時間を要する。

10

【0007】

特許文献1の従来技術は、上記2種類のインデックスを有し、各インデックスを用いた処理に要する時間を見積もり、より短時間で実行可能と判断した処理を実行するものである。しかし、この従来技術は、K to DインデックスとD to Kインデックスの両方を有するため、インデックスを保持するために必要なメモリ領域(記憶容量)が大きい。

【0008】

また、K to Dインデックスに関しては、対象データが追加された場合、追加された対象データおよびその属性に対応させるためにはインデックスを作成しなおさなければならず、更新処理が煩雑であった。

【0009】

そこで本発明は、上記の集計処理を行う技術において、K to DインデックスとD to Kインデックスの両方を用意する場合と比較して、必要なメモリ領域(記憶容量)を削減し、かつ、対象データや属性の更新処理に要する手間を削減することを目的とする。

20

【課題を解決するための手段】

【0010】

上記の目的を達成するため、本発明は、次のようなシステムとして実現される。このシステムは、データを集計するシステムであって、集計処理の対象である対象データの識別情報から対象データが有する属性のリストである第1リストを指定するインデックスであって、対象データごとの第1リストは所定の順序で並び、第1リストの各要素が属性の情報を含む、インデックスと、複数の対象データが有する属性のリストである第2リストと、を記憶するインデックス記憶部と、属性ごとにこの属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計する集計処理を行う集計処理部と、を備える。そして、第2リストの各属性に関して、対象データごとの第1リストの要素を順に辿るリンクが設定されており、集計処理部は、属性ごとに設定されたリンクに基づき、属性を有する対象データを調べる。

30

【0011】

より好ましくは、第1リストにおけるリンクは、対象データごとの第1リストの各要素が、この要素に含まれるのと同じ属性の情報を含む最も近い後続の第1リストの要素を特定する情報を含むことにより、設定される。

【0012】

さらに好ましくは、第2リストの各属性は、複数の対象データにおける個々の対象データが属性を有する頻度の降順にソートされている。

40

【0013】

さらに好ましくは、対象データごとの第1リストの各要素は、この要素が含む情報により特定される属性に対して付与された値をさらに含む。そして、第1リストの各要素は、この値に基づく順序でソートされている。

【0014】

より詳細には、属性に対して付与された値は、この属性とこの属性を有する対象データとの関係に基づく重み値であり、第1リストの各要素は、この値の降順にソートされている。

50

【 0 0 1 5 】

さらに好ましくは、対象データごとの第 1 リストは、最も古い第 1 リストを最後尾として作成された順に並び、新規に作成された第 1 リストを追加する場合は、第 1 リストの並びの先頭に配置され、かつ、新規に作成された第 1 リストの各要素に含まれる属性に関するリンクが更新される。

【 0 0 1 6 】

さらに好ましくは、対象データごとの第 1 リストは、最も古い第 1 リストを最後尾として作成された順に並び、最も古い第 1 リストを削除する場合は、最後尾の第 1 リストが削除され、かつ、削除された第 1 リストの各要素に含まれる属性に関するリンクが更新される。

10

【 0 0 1 7 】

また、本発明は、次のようなシステムとしても実現される。このシステムは、データを集計するシステムであって、集計処理の対象である対象データごとに対象データが有する属性を登録したリストを含むインデックスを記憶するインデックス記憶部と、属性ごとにこの属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計する第 1 手法、および、対象データごとに対象データが有する属性を調べ、対象データとの関係が所定の基準を満たす属性を集計する第 2 手法のどちらで集計処理を行うかを判定する判定部と、判定部による判定に応じて第 1 手法および第 2 手法の一方により集計処理を行う集計処理部と、を備える。そして、インデックスは、相異なる対象データが有する同一の属性どうしを関連付けており、集計処理部は、第 1 手法により集計処理を行う場合に、インデックスにおける同一の属性どうしの関連に基づき、この属性を有する対象データを調べる。

20

【 0 0 1 8 】

また、本発明は、次のような方法としても実現される。この方法は、データを集計する方法であって、集計処理の対象である対象データの識別情報から対象データが有する属性のリストである第 1 リストを指定するインデックスであって、対象データごとの第 1 リストは所定の順序で並び、第 1 リストの各要素がこの属性の情報を含む、インデックスと、複数の対象データが有する属性のリストである第 2 リストと、第 2 リストの各属性に関して、対象データごとの第 1 リストの要素を順に辿るリンク構造と、を記憶するインデックス記憶部と、属性の集計処理を行う集計処理部と、を備えるシステムにおいて、集計処理部が、属性ごとに設定されたリンクに基づき、この属性を有する対象データの数を調べるステップと、属性を、調べた対象データの数が多い上位の所定数の属性を集計し、集計結果を出力するステップと、を含む。

30

【 0 0 1 9 】

また、本発明は、次のような方法としても実現される。この方法は、データを集計する方法であって、集計処理の対象である対象データごとにこの対象データが有する属性を登録したリストを含むインデックスを記憶するインデックス記憶部と、属性の集計処理を行う集計処理部と、この集計処理部により行われる集計処理の手法を判定する判定部と、を備えるシステムにおいて、判定部が、属性ごとにこの属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計する第 1 手法、および、対象データごとにこの対象データが有する属性を調べ、対象データとの関係が所定の基準を満たす属性を集計する第 2 手法のどちらで集計処理を行うかを判定するステップと、判定部が第 1 手法により集計処理を行うと判定した場合に、集計処理部が、相異なる対象データが有する同一の属性どうしを関連付けてなるインデックスにおける同一の属性どうしの関連に基づき、この属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計するステップと、判定部が第 2 手法により集計処理を行うと判定した場合に、集計処理部が、インデックスに基づき、各対象データが有する属性を調べ、対象データとの関係が所定の基準を満たす属性を集計するステップと、を含む。

40

【 0 0 2 0 】

さらにまた、本発明は、コンピュータを制御して上述した装置の各機能を実現するプロ

50

グラム、あるいは、コンピュータに上記の各ステップに対応する処理を実行させるプログラムとしても実現される。このプログラムは、磁気ディスクや光ディスク、半導体メモリ、その他の記録媒体に格納して配布したり、ネットワークを介して配信したりすることにより、提供することができる。

【発明の効果】

【0021】

本発明によれば、K t o D インデックスを別途作成することなく、集計対象の属性ごとに、その属性を有する対象データの数を調べ、対象データ数の多い上位の所定数の属性を集計する。これにより、K t o D インデックスとD t o K インデックスの両方を用意する場合と比較して、必要なメモリ領域（記憶容量）を削減し、かつ、対象データや属性の更新処理に要する手間を削減することが可能となる。

10

【図面の簡単な説明】

【0022】

【図1】本実施形態による集計システムの構成を示す図である。

【図2】本実施形態のインデックス記憶部に記憶されるインデックスの構成を示す図である。

【図3】第1の手法により集計処理を行う場合の集計処理部の動作を示すフローチャートである。

【図4】第1の手法により集計処理を行う場合の集計処理部の動作を示すフローチャートである。

20

【図5】D t o K インデックスの作成処理の手順を示すフローチャートである。

【図6】新規文書追加処理の手順を示すフローチャートである。

【図7】J a v a プログラミング言語を用いたD t o K インデックスの実装例を示す図であり、図7（a）はD t o K インデックスのエントリを設定するコード、図7（b）はD t o K リストを生成するコード、図7（c）は新しいD t o K インデックスを追加するコード、図7（d）はD t o K インデックスを削除するコードを示す図である。

【図8】J a v a プログラミング言語を用いたD t o K インデックスの実装例を示す図であり、図7に示すコードにより設定されたリンクの様子を示す図である。

【図9】J a v a プログラミング言語を用いたD t o K インデックスの他の実装例を示す図であり、D t o K インデックスのエントリを設定するコードを示す図である。

30

【図10】J a v a プログラミング言語を用いたD t o K インデックスの他の実装例を示す図であり、図9に示すコードにより設定されたリンクの様子を示す図である。

【図11】図9～図10に示す実装例によるインデックスの構成を示す図である。

【図12】図9～図11の実装例において、D t o K インデックスを削除する処理の手順を示すフローチャートである。

【図13】図11に示したインデックス構成において、D t o K インデックスを削除した様子を示す図である。

【図14】本実施形態の集計システムを構成するのに好適なハードウェア構成例を示す図である。

【発明を実施するための形態】

40

【0023】

以下、添付図面を参照して、本発明の実施形態について詳細に説明する。本実施形態による集計システムは、複数種類の属性を持つ対象データの集合中に出現する属性を、何らかの規則に基づいて所定の基準個数だけ求める集計処理を行う。以下では、対象データの一例として、検索等の何らかの手段で処理対象として絞り込まれた文書、属性の一例として、各文書に含まれる単語を用いた場合を例として説明する。また、この例において、集計の規則は、絞り込まれた文書全体における出現数の多いものから順に基準個数（k 個）の単語を選択するものとする。

【0024】

< システム構成 >

50

図 1 は、本実施形態による集計システムの構成を示す図である。図 1 に示すように、本実施形態の集計システム 100 は、文書 DB (データベース) 200 に接続されている。文書 DB 200 は、処理対象となり得る文書を格納している。集計処理の対象となる文書は、例えば、文書 DB 200 に格納されている文書のうちで、何らかの検索条件に基づく検索等により指定された (絞り込まれた) 文書とすることができる。

【0025】

本実施形態では、集計システム 100 は、文書 DB 200 に格納されている文書に関して、何らかの検索条件に基づく検索により集計処理の対象となる文書を指定する (絞り込む)。例えば、特定の単語を含む文書を検索し、得られた文書を集計処理の対象とする。そして、得られた文書に含まれる単語のうち、出現頻度の高いものから順に、予め定められた個数 (k 個) の単語を求める。図 1 に示すように、本実施形態の集計システム 100 は、入力受け付け部 110 と、インデックス記憶部 120 と、集計処理部 130 と、判定部 140 と、インデックス更新部 150 とを備える。

【0026】

入力受け付け部 110 は、集計処理を行うための入力を受け付ける。具体的には、入力受け付け部 110 は、集計処理の対象となる文書を指定するための絞り込み条件 (検索条件等) の入力を受け付ける。また、入力受け付け部 110 は、集計しようとする属性を指定する情報の入力を受け付ける。

【0027】

インデックス記憶部 120 は、集計処理に用いられるインデックスを記憶している。本実施形態では、文書の識別情報からその文書に含まれる単語のリストを指定する D t o K インデックスが用いられる。D t o K インデックスは、文書 DB 200 に格納されている文書ごとに作成されている。また、本実施形態では、D t o K インデックスにおける文書ごとの単語のリストにおいて、各文書に出現する同一の単語間にリンクを設定している。言い換えれば、相異なる文書の D t o K インデックスに含まれる同一の単語どうしを関連付けている。これにより、例えば、ある単語をキーワードとして、この単語間のリンクを辿ることにより、キーワードを含む文書を特定することができる。この D t o K インデックスのデータ構造の詳細については後述する。

【0028】

集計処理部 130 は、文書 DB 200 に格納されている文書のうち、集計処理の対象として指定された文書 (以下、指定文書と呼ぶ) を対象として、指定文書全体における出現頻度の高い単語を特定し、集計する。具体的には、集計処理部 130 は、出現頻度が予め定められた上位 k 個に含まれる単語を集計する。本実施形態の集計処理部 130 は、次の何れかの手法を用いて集計処理を行う。

【0029】

すなわち、一つは、集計対象 (集計される属性) の単語ごとに、その単語を含む文書の数を調べ、最も多くの文書に含まれる単語から降順に k 個の単語を選択する処理である。以下、この手法を第 1 手法と呼ぶ。より一般的に言い換えれば、第 1 手法は、属性ごとに、その属性を有する対象データを調べ、対象データとの関係が所定の基準を満たす属性を集計する手法である。

【0030】

他の一つは、指定文書ごとに、その指定文書に含まれる単語を調べ、最も多くの指定文書に含まれる単語から降順に k 個の単語を選択する処理である。以下、この手法を第 2 手法と呼ぶ。より一般的に言い換えれば、第 2 手法は、対象データごとに、その対象データが有する属性を調べ、対象データとの関係が所定の基準を満たす属性を集計する手法である。

【0031】

本実施形態において、第 1 手法の処理は、D t o K インデックスの単語間に設けられたリンクを辿ることによって行われる。集計処理部 130 の動作の詳細については後述する。

10

20

30

40

50

【 0 0 3 2 】

また、集計処理部 1 3 0 は、出力候補保持部 1 3 1 を有している。出力候補保持部 1 3 1 には、集計処理部 1 3 0 による集計処理の過程で取得される、集計結果として出力する単語の候補が保持される。すなわち、集計処理部 1 3 0 は、第 1 手法または第 2 手法による処理において、集計対象となる個々の単語に対して、集計結果として出力する単語に該当するか否かを順次判断する。そして、集計結果として出力する単語に該当すると判断した単語を、出力候補保持部 1 3 1 に保持させる。言い換えれば、出力候補保持部 1 3 1 に保持される単語は、集計処理部 1 3 0 による集計処理の実行途中における、既に処理の済んだ単語である。集計処理が完了した後、集計処理部 1 3 0 は、出力候補保持部 1 3 1 に保持されている単語を、集計結果として出力する。

10

【 0 0 3 3 】

判定部 1 4 0 は、集計処理部 1 3 0 による集計処理を第 1 手法と第 2 手法のどちらで行うかを判定する。上記の第 1 手法による処理は、集計対象である単語の種類が多い場合、多大な処理時間を要する。一方、第 2 手法による処理は、集計処理の対象である文書の数が多い場合、多大な処理時間を要する。そこで、判定部 1 4 0 は、集計処理部 1 3 0 による集計処理が効率良く実行されるように、第 1 手法と第 2 手法のどちらで集計処理を行うかを決定する。なお、具体的な決定方法は、特に限定しない。例えば、集計処理の対象である文書の数が予め設定された閾値よりも多い場合に第 1 手法により集計処理を実行すると決定し、文書の数が閾値以下である場合に第 2 手法により集計処理を実行すると決定しても良い。また、第 1 手法で集計処理を行う場合に要する処理時間と、第 2 手法で集計処理を行う場合に要する処理時間とを予想し、予想される処理時間の短い方の手法により集計処理を行うと決定しても良い。

20

【 0 0 3 4 】

インデックス更新部 1 5 0 は、インデックス記憶部 1 2 0 に記憶されているインデックスを更新する。具体的には、インデックス更新部 1 5 0 は、文書 DB 2 0 0 の更新（文書の追加または削除）に応じて、追加された文書に対応する D t o K インデックスを追加したり、削除したりする。また、インデックス更新部 1 5 0 は、D t o K インデックスの更新に伴い、D t o K インデックスを管理するための D t o K リスト（後述）も更新する。また、インデックス更新部 1 5 0 は、新たに D t o K インデックスを追加する更新処理を繰り返すことにより、インデックス記憶部 1 2 0 に記憶される D t o K インデックス群を作成する。インデックス更新部 1 5 0 による更新処理の詳細については後述する。

30

【 0 0 3 5 】

< インデックスの構成 >

図 2 は、本実施形態のインデックス記憶部 1 2 0 に記憶されるインデックスの構成例を示す図である。図 2 に示すように、本実施形態では、単語リスト 1 2 1 と、文書 DB 2 0 0 に格納されている文書ごとの D t o K インデックス 1 2 2 とが用意される。言い換えれば、この D t o K インデックス 1 2 2 は、文書（対象データ）ごとに作成された、文書に含まれる単語（対象データが有する属性）のリスト（第 1 リスト）である。なお、図 2 に示す例では、n 個の文書に対する D t o K インデックス 1 2 2 が存在し、最後に文書 DB 2 0 0 に格納された文書（最も新しい文書）から順に、(1) ~ (n) の枝番を付している。すなわち、符号 1 2 2 (n) は、n 個の文書の中で最初に文書 DB 2 0 0 に格納された文書（最も古い文書）に対する D t o K インデックス 1 2 2 であり、符号 1 2 2 (1) は、最後に文書 DB 2 0 0 に格納された文書に対する D t o K インデックス 1 2 2 である。すなわち、D t o K インデックス 1 2 2 は、最初に作成された D t o K インデックス 1 2 2 を最後尾として、作成された順に並ぶ。

40

【 0 0 3 6 】

また、特に図示しないが、インデックス記憶部 1 2 0 は、D t o K インデックス 1 2 2 を図 2 に示す枝番の順番で管理するための D t o K リストを記憶している。この D t o K リストは、D t o K インデックス 1 2 2 の識別情報 d i (i : 1 i n) を添え字「 i 」の昇順に並べて登録している。したがって、D t o K リストにおいて、先頭の D t o K

50

インデックス122の識別情報はd1であり、最後のD t o Kインデックス122の識別情報はd nである。また、識別情報d iの添え字「i」は、図2に示したD t o Kインデックス122の枝番(1)~(n)に対応している。集計処理部130は、D t o Kリストから各D t o Kインデックス122(1)~122(n)ヘランダム・アクセスが可能である。また、D t o Kリストは、リストの先頭と終端に対して、エントリ(要素)の追加および削除が可能なリスト構造である。このようなリストは、例えば、J a v a(登録商標)のjava.util.LinkedListにより実現される。

【0037】

単語リスト121は、文書DB200に格納されている文書に含まれる全ての単語の識別情報のリスト(第2リスト)である。この単語リスト121には、各単語に関して、その単語の識別情報、その単語を含む文書の数(以下、出現文書数と呼ぶ)、その単語に関するリンク先の情報が登録される。リンク先の情報については後述する。また、単語リスト121は、各単語の識別情報に関してランダム・アクセスが可能であり、かつ、出現文書数に関してソート順を保持するマップ構造である。このようなリスト(マップ)は、例えば、J a v aのjava.util.TreeMapにより実現される。

【0038】

図2に示す例では、単語リスト121には、m個(識別情報「1」~「m」)の単語が収録されている。単語リスト121に収録された単語は、最も多くの文書に含まれる単語から降順に並べられている。言い換えれば、単語リスト121において、対象データの属性である各単語は、対象データである文書全体における個々の文書に出現する頻度の降順にソートされている。なお、文書DB200が更新されて文書が追加された場合、追加された文書にのみ含まれる単語が存在するならば、その単語に関するエントリが、単語リストの末尾に追加される。すなわち、図2に示す単語リスト121においては、識別情報「m」のエントリの後に、識別情報「m+1」、「m+2」、...、というようにエントリが追加されていく。

【0039】

D t o Kインデックス122は、文書DB200に格納されている文書ごとに作成された、個々の文書に含まれる単語の情報を参照するインデックスである。D t o Kインデックス122には、対応する文書に含まれる各単語に関して、その単語の識別情報と、その単語に付与された点数と、その単語に関するリンク先の情報が登録される。リンク先の情報については後述する。単語に付与される点数は、任意に定義し、設定して良い。例えば、T F - I D F(Term Frequency - Inverse Document Frequency)による重みの値などを点数として用いることができる。本実施形態では、各文書における、その単語の出現数を点数とする。また、D t o Kインデックス122に収録される単語は、この点数の値が最も大きい単語から降順に並べられている。

【0040】

本実施形態において、単語リスト121および各D t o Kインデックス122(1)~122(n)に登録されている各単語は、同一の単語どうしの間にリンクが張られている。本実施形態では、単語リスト121に登録されている識別情報を始点として、D t o Kインデックス122の枝番の順(すなわち、新しい文書から古い文書へ向かう順)にしたがってリンクが張られる。言い換えれば、属性である単語の間のリンクは、対象データである文書ごとのD t o Kインデックス122(1)~122(n)の順序に基づく特定の順序で張られる。また、枝番で次のD t o Kインデックス122に同一の単語が無い場合は、枝番がより後のD t o Kインデックス122であって同一の単語を含むもののうち、直近のD t o Kインデックス122に含まれる単語に対してリンクが張られる。

【0041】

具体的に、識別情報「1」の単語(以下、単語「1」と記す)について、図2を参照して設定されるリンクを説明する。単語「1」のリンクは、まず、単語リスト121の単語「1」から、D t o Kインデックス122(1)の単語「1」へ張られている。そして、D t o Kインデックス122(1)の単語「1」から、D t o Kインデックス122(2

10

20

30

40

50

）の単語「１」へ張られている。以下、図示は省略するが、D t o K インデックス 1 2 2 の枝番の順にしたがって、単語「１」間のリンクが張られている。

【 0 0 4 2 】

ただし、ある文書が単語「１」を含まない場合、その文書の D t o K インデックス 1 2 2 を飛ばして、より後方の D t o K インデックス 1 2 2 に対してリンクが設定される。例えば、図示しない D t o K インデックス 1 2 2 (3) に対応する文書に単語「１」が含まれていない場合を考える。この場合、D t o K インデックス 1 2 2 (2) の単語「１」からのリンクは、D t o K インデックス 1 2 2 (3) を飛ばして、D t o K インデックス 1 2 2 (4) の単語「１」へ張られる。同様に、D t o K インデックス 1 2 2 (4) に対応する文書にも単語「１」が含まれていない場合、D t o K インデックス 1 2 2 (2) の単語「１」からのリンクは、D t o K インデックス 1 2 2 (5) の単語「１」へ張られる。

10

【 0 0 4 3 】

図 2 に示す例では、識別情報「２」の単語は、D t o K インデックス 1 2 2 (1) に対応する文書に含まれていない。したがって、識別情報「２」の単語（以下、単語「２」と記す）のリンクは、D t o K インデックス 1 2 2 (1) を飛ばして、単語リスト 1 2 1 の単語「２」から D t o K インデックス 1 2 2 (2) の単語「２」へ張られている。

【 0 0 4 4 】

したがって、本実施形態によるインデックスの単語間のリンクにおいて、D t o K インデックス 1 2 2 (n) に含まれる各単語は、その単語に関するリンクの終端となる。そして、D t o K インデックス 1 2 2 (n) に含まれていない各単語に関しては、D t o K インデックス 1 2 2 (n) よりも前方の D t o K インデックス 1 2 2 (1) ~ 1 2 2 (n - 1) の何れかにリンクの終端が存在する。

20

【 0 0 4 5 】

本実施形態によるインデックスの単語間のリンクは、既存の種々の手段により設定して良い。一例としては、単語リスト 1 2 1 および各 D t o K インデックス 1 2 2 (1) ~ 1 2 2 (n) における各単語のエントリに、リンク先の D t o K インデックス 1 2 2 およびエントリを指定する情報を記述することによって実現される。他の一例としては、単語リスト 1 2 1 および各 D t o K インデックス 1 2 2 (1) ~ 1 2 2 (n) における各単語のエントリに、他の文書の D t o K インデックス 1 2 2 に含まれている同一の単語のエントリへのポインタを記述することによって実現される。

30

【 0 0 4 6 】

< 集計処理部の動作 >

集計処理部 1 3 0 は、集計処理の対象として指定された指定文書全体における出現頻度の高い単語を特定し集計する処理を、第 1 手法または第 2 手法の一方により行う。第 1 手法は、集計対象（集計される属性）の単語ごとに、その単語を含む指定文書の数を調べ、最も多くの指定文書に含まれる単語から降順に k 個の単語を選択する処理である。また、第 2 手法は、指定文書ごとに、その文書に含まれる単語を調べ、最も多くの文書に含まれる単語から降順に k 個の単語を選択する処理である。

【 0 0 4 7 】

ここで、第 2 手法は、従来の D t o K インデックスを用いた既存の集計処理と同様である。本実施形態においても、図 2 を参照して説明した D t o K インデックス 1 2 2 を用いて従来と同様の処理を行うことができる。すなわち、集計処理部 1 3 0 は、まず、指定文書の D t o K インデックス 1 2 2 に基づいて、各指定文書に含まれる単語のリストを作成する。そして、集計処理部 1 3 0 は、作成したリストに基づき、最も多くの文書に含まれる単語から降順に k 個の単語を選択する。

40

【 0 0 4 8 】

一方、第 1 手法は、従来技術における K t o D インデックスを用いた既存の集計処理と同様の考え方による処理である。ただし、本実施形態では、K t o D インデックスを用いず、図 2 を参照して説明した D t o K インデックス 1 2 2 における単語間のリンクを利用して処理を行う。

50

【 0 0 4 9 】

図 3 および図 4 は、第 1 の手法により集計処理を行う場合の集計処理部 1 3 0 の動作を示すフローチャートである。集計処理部 1 3 0 は、単語リスト 1 2 1 に収録された単語に対する処理を、例えば、識別情報「1」の単語から順に処理対象として着目して行う。すなわち、集計処理部 1 3 0 は、まず、未処理の単語が残っているか否かを判断する（S 3 0 1）。そして、未処理の単語がある場合は（S 3 0 1 で Y e s）、集計処理部 1 3 0 は、未処理の単語から処理対象の単語（以下、対象単語と呼ぶ）を選択する（S 3 0 2）。ここで、対象単語として選択される単語は、例えば、未処理の単語のうちで識別情報の値が最も小さい単語である。単語リスト 1 2 1 に収録された単語は、最も多くの文書に含まれる単語から降順に並べられている。したがって、識別情報の値が最も小さい単語は、未

10

【 0 0 5 0 】

次に、集計処理部 1 3 0 は、出力候補保持部 1 3 1 に保持されている処理済みの単語（以下、処理済み単語と呼ぶ）の数が集計数の k 個か否かを判断する（S 3 0 3）。そして、処理済み単語の数が k 個であれば（S 3 0 3 で Y e s）、集計処理部 1 3 0 は、対象単語の出現文書数「g」を取得する（S 3 0 4）。この出現文書数「g」は、文書 DB 2 0 0 に格納されている文書全体のうち、対象単語を含む文書の数である。

【 0 0 5 1 】

次に、集計処理部 1 3 0 は、対象単語の出現文書数「g」と、出力候補保持部 1 3 1 に保持されている各処理済み単語の出現文書数「h 1」とを比較する（S 3 0 5）。ここで、処理済み単語の出現文書数「h 1」は、その処理済み単語を含む指定文書の数である。すなわち、処理済み単語の出現文書数「h 1」は、文書 DB 2 0 0 に格納されている文書全体のうち、その処理済み単語を含み、かつ、集計処理の対象として指定された文書の数である。後述するように、処理済み単語の出現文書数は、出力候補保持部 1 3 1 に保持されている。対象単語の出現文書数「g」が何れかの処理済み単語の出現文書数「h 1」よりも大きい場合（S 3 0 6 で Y e s）、集計処理部 1 3 0 は、この対象単語に関して、D t o K インデックス 1 2 2 のリンクを辿って、この対象単語を含む文書のリストを取得する（S 3 0 7）。

20

【 0 0 5 2 】

次に、集計処理部 1 3 0 は、取得したリストに含まれる文書のうち、絞り込み条件を満足する文書（指定文書）の数を算出する（S 3 0 8）。そして、集計処理部 1 3 0 は、算出された指定文書の数「h 2」と、各処理済み単語の出現文書数「h 1」とを比較する（S 3 0 9）。指定文書数「h 2」が何れかの処理済み単語の出現文書数「h 1」よりも大きい場合（S 3 1 0 で Y e s）、集計処理部 1 3 0 は、対象単語および処理済み単語のうち、出現文書数の多い方から k 個を選択する。そして、集計処理部 1 3 0 は、選択した各単語と、その単語の出現文書数とを対応付けて、出力候補保持部 1 3 1 に記憶させる。これにより、集計結果として出力する単語の候補である処理済み単語が更新される（S 3 1 1）。なお、処理済み単語の数の上限は k 個なので、今回の処理の対象単語が処理済み単語に追加されたことに伴い、更新前の処理済み単語のうち、最も出現文書数の少ない単語の情報が出力候補保持部 1 3 1 から消去される。この後、集計処理部 1 3 0 の処理は、S 3 0 1 に戻る。そして、未処理の単語が残っていれば（S 3 0 1 で Y e s）、S 3 0 2 以降の処理が行われる。

30

40

【 0 0 5 3 】

全ての処理済み単語の出現文書数「h 1」が指定文書数「h 2」以上であった場合（S 3 1 0 で N o）、今回の処理の対象単語は、処理済み単語に追加されない。したがって、出力候補保持部 1 3 1 に記憶された処理済み単語は更新されることなく、集計処理部 1 3 0 の処理は、S 3 0 1 に戻る。そして、未処理の単語が残っていれば（S 3 0 1 で Y e s）、S 3 0 2 以降の処理が行われる。

【 0 0 5 4 】

ここで、S 3 0 3 において、処理済み単語の数が集計数 k 個に達していない場合を考え

50

る (S 3 0 3 で N o) 。この場合、今回の処理の対象単語は、必ず、集計結果として出力する単語の候補となる。したがって、集計処理部 1 3 0 は、この対象単語に関して、D t o K インデックス 1 2 2 のリンクを辿って、この対象単語を含む文書のリストを取得する (S 3 1 3) 。そして、集計処理部 1 3 0 は、取得したリストに含まれる文書のうち、絞り込み条件を満足する文書 (指定文書) の数を算出する (S 3 1 4) 。この後、集計処理部 1 3 0 は、対象単語と算出した指定文書の数とを対応付けて、出力候補保持部 1 3 1 に記憶させる。これにより、集計結果として出力する単語の候補である処理済み単語が更新される (S 3 1 1) 。なお、この場合は、更新後の処理済み単語の数は k 個以下なので、何れの処理済み単語も、出力候補保持部 1 3 1 から消去されない。この後、集計処理部 1 3 0 の処理は、S 3 0 1 に戻る。そして、未処理の単語が残っていれば (S 3 0 1 で Y e s) 、S 3 0 2 以降の処理が行われる。

10

【 0 0 5 5 】

一方、S 3 0 1 において、未処理の単語が残っていなければ (S 3 0 1 で N o) 、全ての単語に対して上記の処理が済んだので、集計処理部 1 3 0 は、出力候補保持部 1 3 1 に記憶されている k 個の処理済み単語を、集計結果として出力し (S 3 1 2) 、処理を終了する。

【 0 0 5 6 】

また、S 3 0 6 において、全ての処理済み単語の出現文書数「h 1」が対象単語の出現文書数「g」以上であった場合を考える (S 3 0 6 で N o) 。この場合、単語リスト 1 2 1 に収録された単語は、最も多くの文書に含まれる単語から降順に並べられているため、今回よりも後の処理において対象単語となる単語の出現文書数が何れかの処理済み単語の出現文書数「h 1」よりも大きくなることはない。すなわち、今回よりも後の処理において処理済み単語が更新されることはない。したがって、集計処理部 1 3 0 は、出力候補保持部 1 3 1 に記憶されている k 個の処理済み単語を、集計結果として出力し (S 3 1 2) 、処理を終了する。

20

【 0 0 5 7 】

このように、S 3 0 5、S 3 0 6 に示す処理によれば、単語リスト 1 2 1 に収録された全ての単語に対して処理を行っていない段階でも、処理済み単語が更新される可能性がなくなった場合に、集計処理を終了する。これにより、S 3 0 1 ~ S 3 1 1 の処理の繰り返し回数を、単語リスト 1 2 1 に収録された単語の数よりも少なく抑えることができる。

30

【 0 0 5 8 】

以上のように、本実施形態は、収録した単語間にリンクを設定した D t o K インデックス 1 2 2 を用いて、第 1 手法による処理および第 2 手法による処理のどちらも実行できる。したがって、第 1 手法による処理を行うために K t o D インデックスを用意する場合と比較して、インデックスを保持するために必要なメモリ領域 (記憶容量) を小さく抑えることができる。

【 0 0 5 9 】

なお、単にインデックスを保持するために必要なメモリ領域 (記憶容量) を削減するのであれば、適当な圧縮方式によりインデックスをデータ圧縮してデータサイズを小さくすることが可能である。しかし、文書 D B 2 0 0 の更新 (文書の追加または削除) に応じてデータ圧縮されたインデックスを更新することは困難である。そのため、文書 D B 2 0 0 が更新された場合は、インデックス自体を作成し直し、改めてデータ圧縮することが必要となる。これに対し、本実施形態では、限定的ではあるが、後述のように、文書 D B 2 0 0 の更新に応じて D t o K インデックス 1 2 2 を更新可能としながら、上記のようにインデックスを保持するために必要なメモリ領域 (記憶容量) の削減を実現している。

40

【 0 0 6 0 】

< D t o K インデックスの作成および更新 >

本実施形態において、文書 D B 2 0 0 が更新されると、これに伴って、インデックス記憶部 1 2 0 に記憶された単語リスト 1 2 1 および D t o K インデックス 1 2 2 も更新される。単語リスト 1 2 1 に関しては、文書の追加または削除により、収録している各単語を

50

含む文書の数が変わる。そこで、本実施形態の集計システム 100 のインデックス更新部 150 は、例えば、文書 DB 200 が更新される度に、単語リスト 121 に収録している単語をソートし直す。また、本実施形態のインデックス更新部 150 は、追加された文書にのみ含まれる単語が存在するならば、その単語の識別情報を、単語リストの末尾に追加することにより、単語リスト 121 を更新する。

【0061】

次に、D t o K インデックス 122 の作成および更新について説明する。D t o K インデックス 122 の作成は、一つの文書に対応する新たな D t o K インデックス 122 を追加する更新処理を、文書 DB 200 に格納された n 個の文書に対して実行することにより行われる。

10

【0062】

図 5 は、D t o K インデックス 122 の作成処理の手順を示すフローチャートである。図 5 に示すように、インデックス更新部 150 は、まず、文書 DB 200 に格納されている文書のうち、未処理の (D t o K インデックス 122 の作成が行われていない) 文書があるか否かを調べる (S 401)。未処理の文書があれば (S 401 で Y e s)、インデックス更新部 150 は、未処理の文書の一つを取得する (S 402)。そして、インデックス更新部 150 は、取得した文書を処理対象として、新規文書追加処理を行う (S 403)。以上の処理を文書 DB 200 に格納されている各文書に対して行い、未処理の文書が無くなったならば (S 401 で N o)、インデックス更新部 150 は、処理を終了する。

20

【0063】

図 6 は、新規文書追加処理の手順を示すフローチャートである。図 6 に示すように、インデックス更新部 150 は、まず、処理対象として取得した一つの文書に対応する D t o K インデックス 122 を作成する (S 501)。以下、作成された D t o K インデックス 122 を、D t o K リストに登録される各 D t o K インデックス 122 の識別情報 d i を用いてインデックス d i と記載する (図 6 の S 501 でもインデックス d i と記載)。また、インデックス d i に対応する文書を文書 d i と記載する。インデックス更新部 150 は、作成したインデックス d i を D t o K リストの先頭に追加する (S 502)。

【0064】

次に、インデックス更新部 150 は、インデックス d i に登録されている単語 (文書 d i に含まれる単語) のうち、未処理の (リンクの設定が更新されていない) 単語があるか否かを調べる (S 503)。未処理の単語がある場合 (S 503 で Y e s)、インデックス更新部 150 は、未処理の単語の一つを選択し (S 504)、選択した単語が単語リスト 121 に登録されているか否かを調べる (S 505)。そして、選択した単語が登録されているならば (S 505 で Y e s)、インデックス更新部 150 は、単語リスト 121 におけるその単語のエントリに登録されているリンク先の情報を、処理中のインデックス d i におけるその単語のエントリを指すように更新する (S 507)。

30

【0065】

一方、選択した単語が単語リスト 121 に登録されていない場合 (S 505 で N o)、インデックス更新部 150 は、選択した単語のエントリを単語リスト 121 に追加する (S 506)。そして、インデックス更新部 150 は、単語リスト 121 において作成したエントリのリンク先の情報を、処理中のインデックス d i におけるその単語のエントリを指すように更新する (S 507)。

40

【0066】

次に、インデックス更新部 150 は、インデックス d i におけるその単語のエントリに登録されているリンク先の情報を、D t o K リストにおいて後方かつ直近の他の D t o K インデックス 122 におけるその単語のエントリを指すように更新する (S 508)。

【0067】

インデックス更新部 150 は、以上の S 503 ~ S 508 の処理を、インデックス d i に含まれる各単語に対して行い、未処理の単語が無くなったならば (S 503 で N o)、

50

インデックス更新部 150 は、処理を終了する。

【0068】

以上、D t o K インデックス 122 を追加する処理について説明した。次に、インデックス記憶部 120 に記憶されている D t o K インデックス 122 を削除する処理について説明する。図 2 を参照して説明した本実施形態の D t o K インデックス 122 は、単語間にリンクが設定されている。そのため、D t o K インデックス 122 を削除する場合は、削除対象の D t o K インデックス 122 における単語のリンクを設定し直す必要がある。したがって、任意の D t o K インデックス 122 を削除することは容易ではない。しかし、D t o K リストの終端に対応する（すなわち、最も古い）D t o K インデックス 122（識別情報 d n）については、削除することが容易である。これは、D t o K インデックス 122 の全ての単語が、その単語に関するリンクの終端だからである。この場合、インデックス記憶部 120 から識別情報 d n の D t o K インデックス 122（図 2 に示す例では、D t o K インデックス 122（n））を削除し、D t o K リストから終端のエントリを削除すれば良い。

10

【0069】

< インデックスの実装例 >

図 7 および図 8 は、J a v a プログラミング言語を用いた、本実施形態による D t o K インデックス 122 の実装例を示す図である。図 7 は、コードの例を示す図であり、図 7（a）は D t o K インデックス 122 のエントリを設定するコード、図 7（b）は D t o K リストを生成するコード、図 7（c）は新しい D t o K インデックス 122 を追加するコード、図 7（d）は D t o K インデックス 122 を削除するコードである。図 8 は、図 7 に示すコードにより設定されたリンクの様子を示す図である。

20

【0070】

図 7 に示すように実装した場合、D t o K インデックス 122 の単語（エントリ）間のリンクは、リンク先の D t o K インデックス 122 およびエントリの位置を数値で指定することにより実現される。図 7（a）に示す例では、リンク先の D t o K インデックス 122 は、「nextDocId」の値で指定される。ここで、D o c I d とは、後述の文書 I D である。また、リンク先のエントリの位置は、「nextEntryIndex」の値で指定される。また、「keywordId」は、各単語の識別情報である。そして、「score」は、各単語に付与された点数である。また、図 7（b）に示すように、J a v a の標準ライブラリに存在するクラスである「java.util.TreeMap」により、D t o K リスト「DLList」が生成される。この D t o K リスト「DLList」は、D t o K インデックス 122（図 7 では D L [i] と記載）の全体を識別情報 d i の昇順（D L [1] ~ D L [n]）に格納する。

30

【0071】

また、図 7（c）に示す例では、D t o K インデックス 122 が追加される場合に、追加される D t o K インデックス 122 に対応する文書の識別情報（文書 I D）が設定される。この文書 I D は、D t o K リスト「DLList」に登録される D t o K インデックス 122 の識別情報 d i とは別に設定される。例えば、それまでに使用した文書 I D の最大値を記憶しておき、その最大値よりも 1 大きい値を新たに設定する文書 I D の値とする。このようにすれば、図 7（d）に示すように、文書 I D が最小値である D t o K インデックス 122 を削除することにより、D t o K リスト「DLList」の終端のエントリに対応する D t o K インデックス 122 が削除される。

40

【0072】

図 8 に示す例では、識別情報「200」の単語に関して、D t o K インデックス 122（3）から D t o K インデックス 122（5）へ設定されたリンクが示されている。図示の例において、D t o K インデックス 122（3）の文書 I D は「10」であり、D t o K インデックス 122（5）の文書 I D は「6」である。また、D t o K インデックス 122（3）において、識別情報「200」の単語は、2 番目のエントリ「Entry[1]」に登録されている。D t o K インデックス 122（3）のエントリ「Entry[1]」を参照すると、識別情報「200」の単語の点数は「1.5」である。また、リンク先は、文書 I D 「

50

6」のD t o Kインデックス1 2 2のエントリ「Entry[2]」となっている。したがって、リンク先である、D t o Kインデックス1 2 2 (5) の3 番目のエントリ「Entry[2]」には、識別情報「2 0 0」の単語が登録されている。

【0 0 7 3】

図8に示す例において、文書ID「6」の文書が削除された場合を考える。D t o Kリスト「DLList」およびD t o Kインデックス1 2 2が更新されると、文書ID「6」のD t o Kインデックス1 2 2 (5) が削除される。これにより、D t o Kリスト「DLList」に文書ID「6」のD t o Kインデックス1 2 2 (5) が存在しなくなる。そのため、D t o Kインデックス1 2 2 (3) のエントリ「Entry[1]」からD t o Kインデックス1 2 2 (5) のエントリ「Entry[2]」へのリンクを辿れなくなる。そして、識別情報「2 0 0」の単語に関するリンクは、D t o Kインデックス1 2 2 (3) のエントリ「Entry[1]」が終端となる。

10

【0 0 7 4】

< インデックスの他の構成および実装例 >

以上の実施形態において、単語リスト1 2 1および各D t o Kインデックス1 2 2 (1) ~ 1 2 2 (n) に登録されている単語間のリンクは、単語リスト1 2 1に登録されている識別情報を始点として、識別情報d iの順にしたがって設定した。これに対し、D t o Kインデックス1 2 2の実装によっては、識別情報d iの降順（すなわち、古い文書から新しい文書へ向かう順）にリンクを設定することが望ましい場合がある。

【0 0 7 5】

20

図9および図10は、J a v aプログラミング言語を用いた、本実施形態によるD t o Kインデックス1 2 2の他の実装例を示す図である。図9は、D t o Kインデックス1 2 2のエントリを設定するコードの例を示す図である。図10は、図9に示すコードにより設定されたリンクの様子を示す図である。なお、D t o Kリストを生成するコード、新しいD t o Kインデックス1 2 2を追加するコード、D t o Kインデックス1 2 2を削除するコードは、図7 (b) ~ (d) に示したコードと同様とする。

【0 0 7 6】

図9に示すように実装した場合、D t o Kインデックス1 2 2の単語（エントリ）間のリンクは、リンク先のエントリをポインタにより直接指し示すことにより実現される。図9に示す例では、リンク先のD t o Kインデックス1 2 2のエントリが、J a v aやC / C + +等のプログラミング言語でサポートされているポインタ（または同等の機能）で指定される。図示の例では、「Entry next」にポインタが記録される。このように、ポインタを用いてリンクが設定される場合、図7および図8に示した実装例とは反対の向きにリンクを設定する。すなわち、ある単語に関するリンクは、その単語が出現する最も古い（識別情報d iの値が最も大きい）D t o Kインデックス1 2 2のエントリを始点とし、単語リスト1 2 1におけるその単語のエントリを終端とする。

30

【0 0 7 7】

また、この実装例では、単語リスト1 2 1のエントリから、各単語が出現する最も古いD t o Kインデックス1 2 2のエントリへ向かうリンク（以下、逆向きリンクと呼ぶ）も設定される。この逆向きリンクは、集計処理部1 3 0が第1手法による集計処理を行う場合に、単語のリンクを辿るために用いられる。すなわち、集計処理部1 3 0は、逆向きリンクにより、通常のリンクの始点を見つける。

40

【0 0 7 8】

図10に示す例では、識別情報「2 0 0」の単語に関して、D t o Kインデックス1 2 2 (5) からD t o Kインデックス1 2 2 (3) へ、図8に示した例とは反対の向きに設定されたリンクが示されている。図示の例において、D t o Kインデックス1 2 2 (3) の文書IDは「1 0」であり、D t o Kインデックス1 2 2 (5) の文書IDは「6」である。また、D t o Kインデックス1 2 2 (5) において、識別情報「2 0 0」の単語は、3 番目のエントリ「Entry[2]」に登録されている。D t o Kインデックス1 2 2 (5) のエントリ「Entry[2]」を参照すると、識別情報「2 0 0」の単語の点数は「1 . 5」で

50

ある。そして、リンク先（D t o K インデックス 1 2 2 (3) のエントリ「Entry[1]」）へのポインタが設定されている。したがって、リンク先である、D t o K インデックス 1 2 2 (3) のエントリ「Entry[1]」には、識別情報「2 0 0」の単語が登録されている。
【0 0 7 9】

図 1 1 は、図 9 および図 1 0 に示す実装例によるインデックスの構成を示す図である。図 1 1 に示す構成例は、図 2 に示した構成例と同様に、単語リスト 1 2 1 と、文書 D B 2 0 0 に格納されている文書ごとの D t o K インデックス 1 2 2 とを含む。そして、n 個の文書に対する D t o K インデックス 1 2 2 が存在し、最後に文書 D B 2 0 0 に格納された文書（最も新しい文書）から順に、(1) ~ (n) の枝番を付している。

【0 0 8 0】

10

図 1 1 に示すインデックス構成では、単語ごとに、その単語が出現する最も古い D t o K インデックス 1 2 2 のエントリから単語リスト 1 2 1 におけるその単語のエントリへ至る通常のリンクと、逆向きリンクとが張られている。図 1 1 においては、通常のリンクを実線で示し、逆向きリンクを破線で示している。

【0 0 8 1】

具体的に、識別情報「1」の単語（単語「1」）について、図 1 1 に示すリンクを説明する。単語「1」の通常のリンクは、まず、単語「1」が出現する最も古い D t o K インデックス 1 2 2 (n) のエントリから、D t o K インデックス 1 2 2 (n - 1) のエントリへ張られている。以下、図示は省略するが、単語「1」間のリンクは、D t o K インデックス 1 2 2 の枝番の降順にしたがって張られ、D t o K インデックス 1 2 2 (2) のエントリに至っている。そして、単語「1」のリンクは、D t o K インデックス 1 2 2 (2) のエントリから、D t o K インデックス 1 2 2 (1) のエントリを経て、単語リスト 1 2 1 における単語「1」のエントリを終端としている。

20

【0 0 8 2】

なお、単語「1」を含まない文書の D t o K インデックス 1 2 2 を飛ばしてリンクが張られることは、図 2 に示した構成例と同様である。例えば、図 1 1 において、単語「2」のリンクは、D t o K インデックス 1 2 2 (1) を飛ばして、D t o K インデックス 1 2 2 (2) のエントリから単語リスト 1 2 1 における単語「2」のエントリへ張られている。すなわち、単語「2」は、D t o K インデックス 1 2 2 (1) の文書には出現しない。また、単語「3」のリンクは、D t o K インデックス 1 2 2 (n) のエントリから D t o K インデックス 1 2 2 (2) のエントリへ張られている。すなわち、単語「3」は、図示しない D t o K インデックス 1 2 2 (3) から D t o K インデックス 1 2 2 (n - 2) の文書には出現しない。

30

【0 0 8 3】

また、図 1 1 に示すインデックス構成において、ある単語に関する通常のリンクの始点は、その単語が出現する最も古い D t o K インデックス 1 2 2 であり、必ずしも全体で最も古い D t o K インデックス 1 1 2 (n) ではない。例えば、図 1 1 の例において、単語「2」に関するリンクを考える。単語「2」のリンクの始点は、D t o K インデックス 1 2 2 (n) および D t o K インデックス 1 2 2 (n - 1) のどちらにも存在しない。したがって、単語「2」のリンクの始点は、D t o K インデックス 1 2 2 (n - 1) よりも新しい何れかの D t o K インデックス 1 2 2 である。

40

【0 0 8 4】

また、単語「1」の逆向きリンクは、単語リスト 1 2 1 における単語「1」のエントリから、単語「1」が出現する最も古い D t o K インデックス 1 2 2 (n) のエントリへ張られている。単語「3」の逆向きリンクも同様である。これに対し、単語「2」の逆向きリンクは、単語「2」が出現する最も古い D t o K インデックス 1 2 2 が図示されていないため、記載を省略している。すなわち、単語「2」が出現する最も古い D t o K インデックス 1 2 2 は、D t o K インデックス 1 2 2 (3) ~ D t o K インデックス 1 2 2 (n - 2) の何れかである。

【0 0 8 5】

50

図 9 ~ 図 11 を参照して説明した実装例において、新たな D t o K インデックス 1 2 2 を追加する更新については、リンクの方向が反対方向になることを除き、図 6 に示した手順と同様の手順で実行することができる。リンクの設定に関しては、図 6 の S 5 0 5 で、選択した単語が単語リスト 1 2 1 に登録されているならば、(S 5 0 5 で Y e s)、インデックス更新部 1 5 0 は、追加された D t o K インデックス 1 2 2 におけるその単語のエントリのポインタを、単語リスト 1 2 1 におけるその単語のエントリを指すように更新する。また、インデックス更新部 1 5 0 は、単語リスト 1 2 1 におけるその単語のエントリから逆向きリンクおよび通常のリンクを辿り、現在、単語リスト 1 2 1 におけるその単語のエントリを指すポインタを有する D t o K インデックス 1 2 2 のエントリを検出する。そして、インデックス更新部 1 5 0 は、検出した D t o K インデックス 1 2 2 のエントリのポインタを、追加した D t o K インデックス 1 2 2 におけるその単語のエントリを指すように更新する。

10

【 0 0 8 6 】

一方、図 6 の S 5 0 5 で、選択した単語が単語リスト 1 2 1 に登録されていない場合は (S 5 0 5 で N o)、インデックス更新部 1 5 0 は、選択した単語のエントリを単語リスト 1 2 1 に追加する。そして、インデックス更新部 1 5 0 は、追加した D t o K インデックス 1 2 2 におけるその単語のエントリのポインタを、単語リスト 1 2 1 におけるその単語のエントリを指すように更新してリンクを設定する。また、インデックス更新部 1 5 0 は、単語リスト 1 2 1 におけるその単語のエントリのポインタを、追加した D t o K インデックス 1 2 2 におけるその単語のエントリを指すように更新して逆向きリンクを設定する。

20

【 0 0 8 7 】

これに対し、D t o K インデックス 1 2 2 を削除する更新の場合、単に D t o K インデックス 1 2 2 を削除するだけでなく、削除対象の D t o K インデックス 1 2 2 に収録されている各単語に関して設定されたリンクを更新する。削除対象である最も古い D t o K インデックス 1 2 2 に収録されている各単語は、その単語に関する通常のリンクの始点であり、かつ、その単語に関する逆向きリンクのリンク先である。そのため、D t o K インデックス 1 2 2 を削除した後に、これらのリンクおよび逆向きリンクを正しく辿れるように、リンクを更新することが必要である。

【 0 0 8 8 】

図 12 は、図 9 ~ 図 11 の実装例において、D t o K インデックス 1 2 2 を削除する処理の手順を示すフローチャートである。図 12 に示すように、インデックス更新部 1 5 0 は、まず、削除対象の D t o K インデックス 1 2 2 (n) に登録されている単語のうち、未処理の (リンクの設定が更新されていない) 単語があるか否かを調べる (S 1 1 0 1)。未処理の単語がある場合 (S 1 1 0 1 で Y e s)、インデックス更新部 1 5 0 は、未処理の単語の一つを選択し (S 1 1 0 2)、選択した単語に対して張られている逆向きリンクのリンク先を更新する (S 1 1 0 3)。

30

【 0 0 8 9 】

具体的には、インデックス更新部 1 5 0 は、削除対象の D t o K インデックス 1 2 2 (n) における選択した単語のエントリに登録されている通常のリンクのリンク先へのポインタの情報を取得する。このポインタにより示されるリンク先の D t o K インデックス 1 2 2 は、削除対象の D t o K インデックス 1 2 2 (n) が削除された後に、選択した単語が出現する最も古い D t o K インデックス 1 2 2 となる。そして、インデックス更新部 1 5 0 は、単語リスト 1 2 1 における選択した単語に関する逆向きリンクのリンク先へのポインタを、取得したポインタの情報に基づいて更新する。

40

【 0 0 9 0 】

インデックス更新部 1 5 0 は、以上の S 1 1 0 1 ~ S 1 1 0 3 の処理を、削除対象の D t o K インデックス 1 2 2 (n) に登録されている各単語に対して行う。これにより、D t o K インデックス 1 2 2 (n) の各単語に対して設定されていた逆向きリンクが、D t o K インデックス 1 2 2 (n) の削除後に各単語が出現する最も古い D t o K インデック

50

ス 1 2 2 へ、それぞれ設定し直される。そして、未処理の単語が無くなったならば (S 1 1 0 1 で N o)、インデックス更新部 1 5 0 は、削除対象の D t o K インデックス 1 2 2 (n) を削除して (S 1 1 0 4)、処理を終了する。

【 0 0 9 1 】

図 1 3 は、図 1 1 に示したインデックス構成において、D t o K インデックス 1 2 2 (n) を削除した様子を示す図である。図 1 3 を参照すると、D t o K インデックス 1 2 2 (n) を削除したことにより、単語「 1 」に関する通常のリンクの始点 (逆向きリンクのリンク先) は、D t o K インデックス 1 2 2 (n - 1) における単語「 1 」のエントリとなっている。また、単語「 2 」に関する通常のリンクの始点 (逆向きリンクのリンク先) は、D t o K インデックス 1 2 2 (2) における単語「 2 」のエントリとなっている。

10

【 0 0 9 2 】

< ハードウェア構成例 >

図 1 4 は、本実施形態の集計システム 1 0 0 を構成するのに好適なハードウェア構成例を示す図である。ここでは、コンピュータにより構成する場合について説明する。図 1 4 に示すコンピュータは、演算手段である C P U (Central Processing Unit) 1 0 a と、主記憶手段であるメモリ 1 0 c を備える。また、外部デバイスとして、磁気ディスク装置 (H D D : Hard Disk Drive) 1 0 g、ネットワーク・インターフェイス 1 0 f、ディスプレイ装置を含む表示機構 1 0 d、キーボードやマウス等の入力デバイス 1 0 i 等を備える。

【 0 0 9 3 】

20

図 1 4 に示す構成例では、メモリ 1 0 c および表示機構 1 0 d は、システム・コントローラ 1 0 b を介して C P U 1 0 a に接続されている。また、ネットワーク・インターフェイス 1 0 f、磁気ディスク装置 1 0 g および入力デバイス 1 0 i は、I / O コントローラ 1 0 e を介してシステム・コントローラ 1 0 b と接続されている。各構成要素は、システム・バスや入出力バス等の各種のバスによって接続される。

【 0 0 9 4 】

図 1 4 において、磁気ディスク装置 1 0 g には O S のプログラムやアプリケーション・プログラムが格納されている。そして、これらのプログラムがメモリ 1 0 c に読み込まれて C P U 1 0 a に実行されることにより、集計システム 1 0 0 における集計処理部 1 3 0、判定部 1 4 0 およびインデックス更新部 1 5 0 の各機能が実現される。また、メモリ 1 0 c や磁気ディスク装置 1 0 g により、インデックス記憶部 1 2 0 が実現される。また、入力デバイス 1 0 i およびプログラム制御された C P U 1 0 a により、入力受け付け部 1 1 0 が実現される。さらにまた、本実施形態では、文書 D B 2 0 0 を磁気ディスク装置 1 0 g により実現しても良い。なお、図 1 4 は、本実施形態の集計システム 1 0 0 を実現するのに好適なコンピュータのハードウェア構成を例示するに過ぎず、集計システム 1 0 0 の具体的構成は、図 1 4 に示す構成に限定されない。

30

【 0 0 9 5 】

以上、本実施形態について説明したが、本実施形態は、上記の具体的構成に限定されるものではない。本実施形態は、複数種類の属性を持つ対象データの集合中に出現する属性を、何らかの規則に基づいて各属性に付与された点数の高いものから順に所定個数求める集計処理に対して、広く適用できるものである。すなわち、本実施形態が適用される集計処理における対象データは、上記の処理対象として絞り込まれた文書には限定されない。また、本実施形態が適用される集計処理における属性は、各文書に含まれる単語に限定されない。その他、上記の実施形態に、種々の変更または改良を加えたものも、本発明の技術的範囲に含まれる。

40

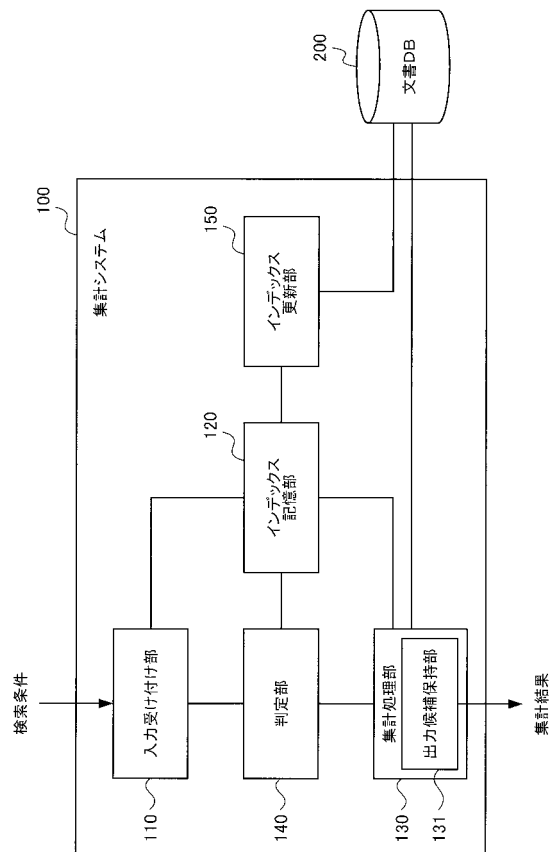
【 符号の説明 】

【 0 0 9 6 】

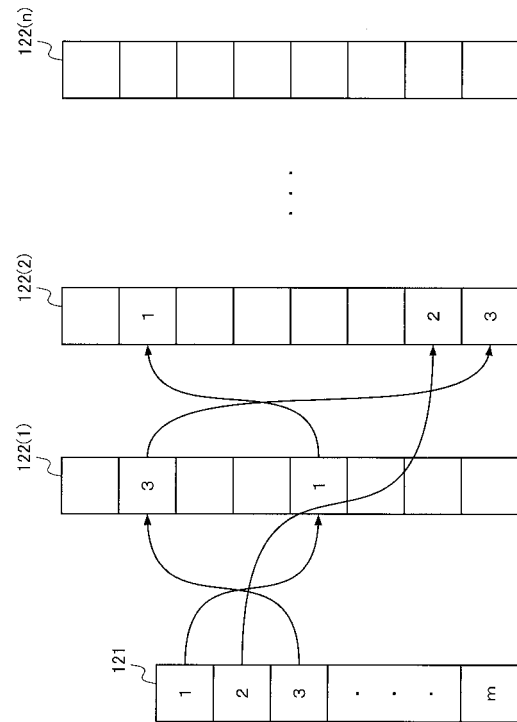
1 0 0 ... 集計システム、 1 1 0 ... 入力受け付け部、 1 2 0 ... インデックス記憶部、 1 2 1 ... 単語リスト、 1 2 2 ... D t o K インデックス、 1 3 0 ... 集計処理部、 1 4 0 ... 判定部、 1 5 0 ... インデックス更新部

50

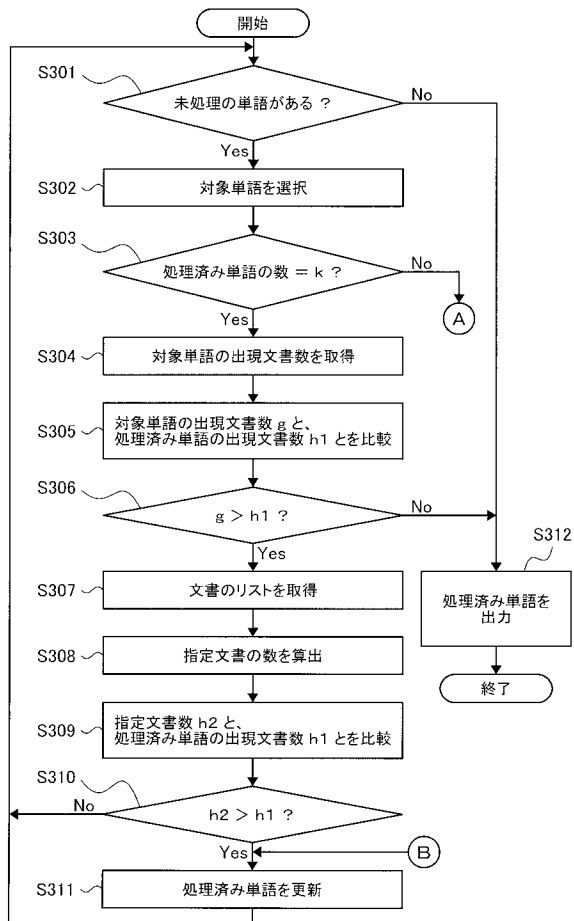
【図 1】



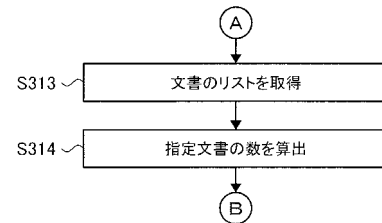
【図 2】



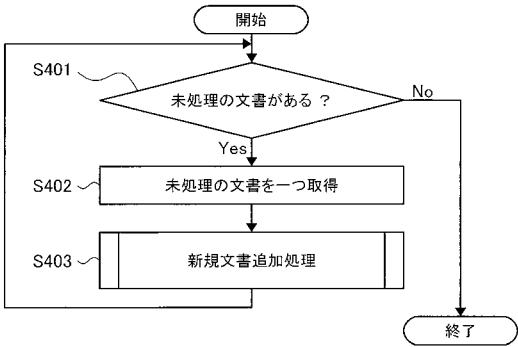
【図 3】



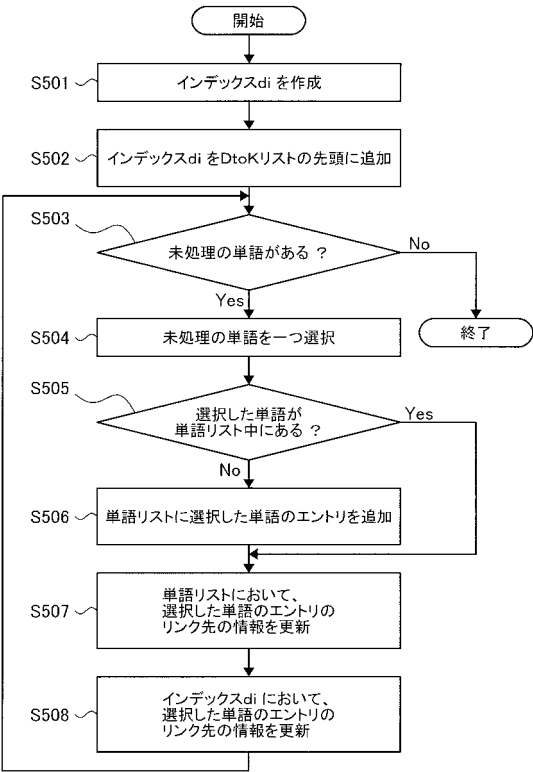
【図 4】



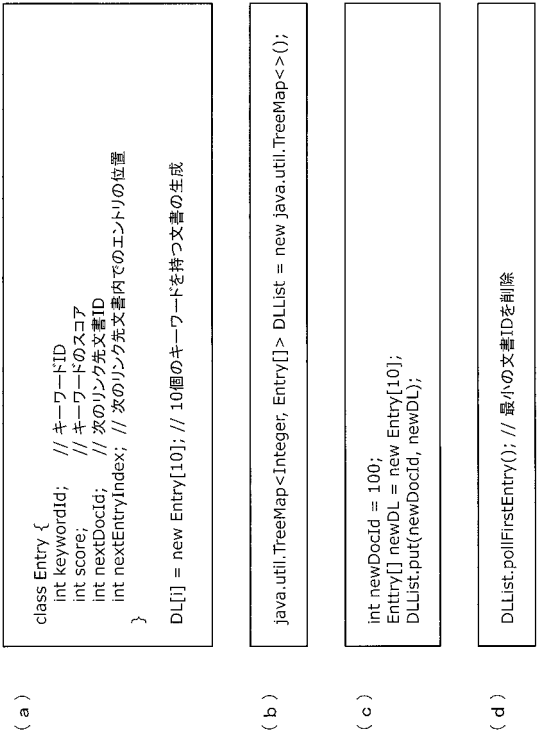
【 図 5 】



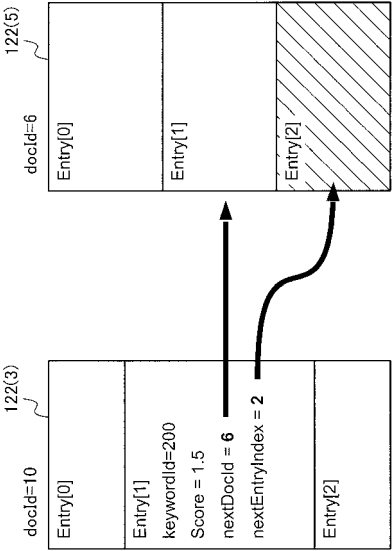
【 図 6 】



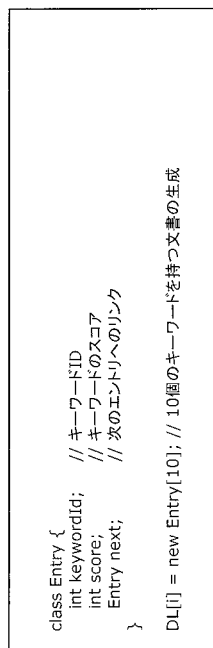
【 図 7 】



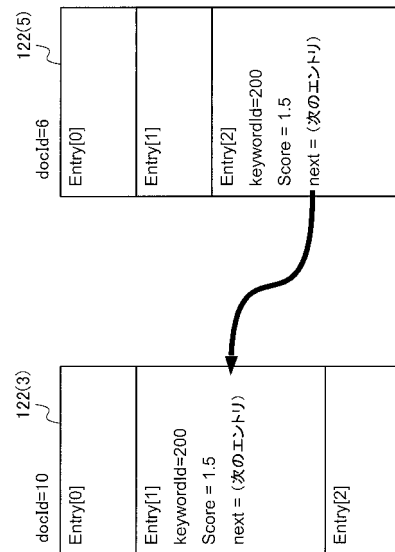
【 図 8 】



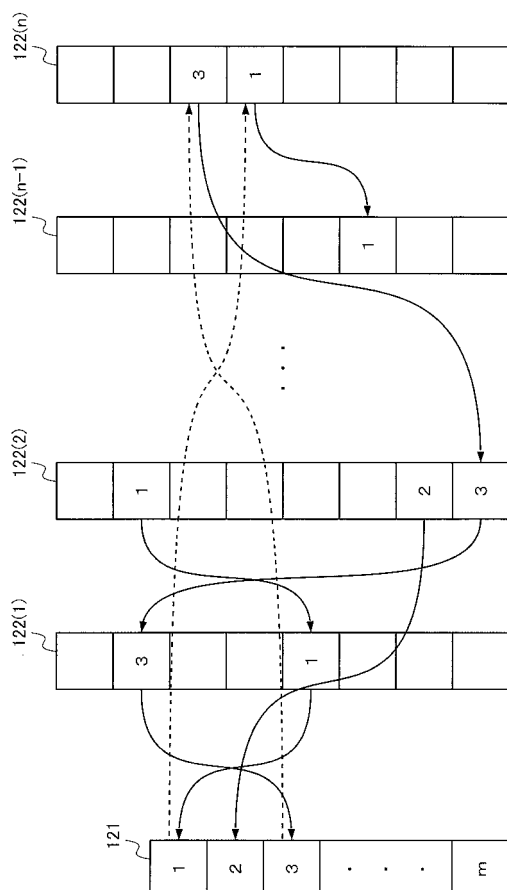
【図 9】



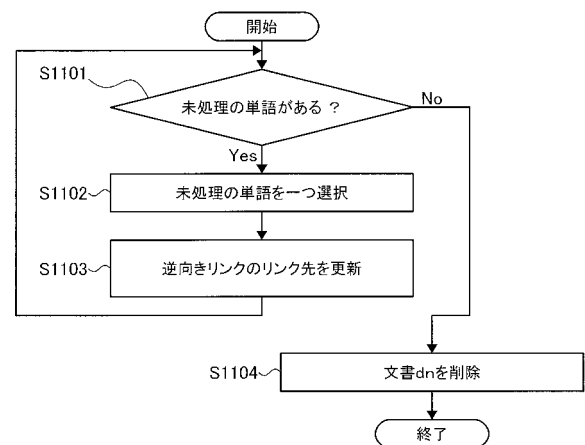
【図 10】



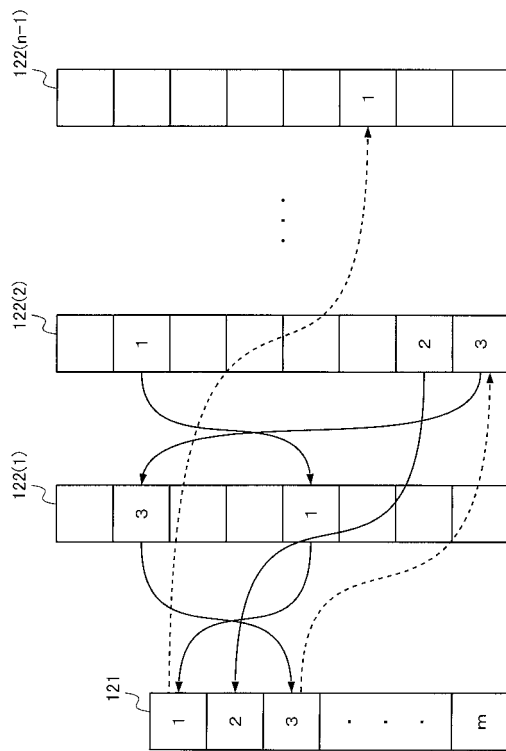
【図 11】



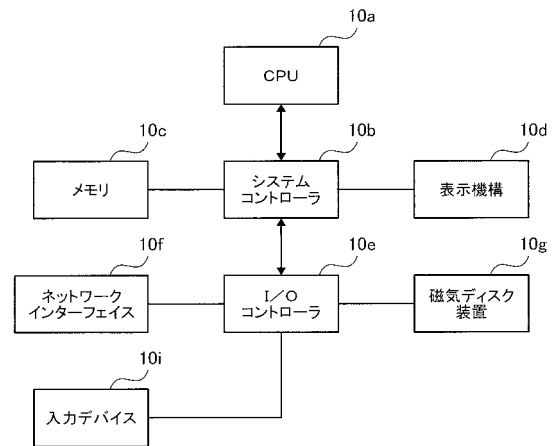
【図 12】



【図 13】



【図 14】



フロントページの続き

(74)代理人 100112690

弁理士 太佐 種一

(74)復代理人 100104880

弁理士 古部 次郎

(74)復代理人 100118108

弁理士 久保 洋之

(72)発明者 吉田 一星

東京都江東区豊洲五丁目6番52号 NBF豊洲キャナルフロント 日本アイ・ビー・エム株式会
社 東京基礎研究所内

(72)発明者 榎 美紀

東京都江東区豊洲五丁目6番52号 NBF豊洲キャナルフロント 日本アイ・ビー・エム株式会
社 東京基礎研究所内