

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第5939123号
(P5939123)

(45) 発行日 平成28年6月22日(2016.6.22)

(24) 登録日 平成28年5月27日(2016.5.27)

(51) Int.Cl.

F I

G 0 6 F 17/30 (2006.01)

G 0 6 F 17/30 1 1 0 C

G 0 6 F 17/30 4 1 5

請求項の数 6 (全 27 頁)

(21) 出願番号 特願2012-224604 (P2012-224604)
 (22) 出願日 平成24年10月9日(2012.10.9)
 (65) 公開番号 特開2014-78085 (P2014-78085A)
 (43) 公開日 平成26年5月1日(2014.5.1)
 審査請求日 平成27年6月4日(2015.6.4)

(73) 特許権者 000005223
 富士通株式会社
 神奈川県川崎市中原区上小田中4丁目1番
 1号
 (74) 代理人 100089118
 弁理士 酒井 宏明
 (72) 発明者 上田 晴康
 神奈川県川崎市中原区上小田中4丁目1番
 1号 富士通株式会社内
 (72) 発明者 前田 高光
 神奈川県川崎市中原区上小田中4丁目1番
 1号 富士通株式会社内

審査官 川▲崎▼ 博章

最終頁に続く

(54) 【発明の名称】 実行制御プログラム、実行制御方法および情報処理装置

(57) 【特許請求の範囲】

【請求項1】

コンピュータに、
 形式の異なる複数の入力ファイルを読み込み、
 各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間フ
 ァイルを、前記入力ファイルごとに生成し、
 前記突合キーに基づいて、各中間ファイル内のデータをソートし、
 データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれ
 ぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを
 生成し、

生成した複数の出力ファイルを、データを突合する突合プログラムに入力する
 処理を実行させることを特徴とする実行制御プログラム。

【請求項2】

前記出力ファイルを生成する処理は、前記各中間ファイルのレコードごと、前記入力フ
 ァイルの各形式のデータを抽出し、

前記入力する処理は、前記各中間ファイルのレコードから前記各形式のデータが抽出さ
 れるたびに、抽出されたデータをプロセス間通信で前記突合プログラムに出力することを
 特徴とする請求項1に記載の実行制御プログラム。

【請求項3】

前記コンピュータは、複数のコンピュータに分散して保持されるデータを分散処理する

分散処理システムを形成し、

前記突合プログラムの処理結果を、他のコンピュータが共有でアクセスできる前記分散処理システム上の共有ファイルシステムに格納する処理をさらにコンピュータに実行させることを特徴とする請求項 1 または 2 に記載の実行制御プログラム。

【請求項 4】

前記格納する処理は、前記突合プログラムが処理結果を出力するたびに、プロセス間通信で、前記共有ファイルシステムに処理結果を格納することを特徴とする請求項 3 に記載の実行制御プログラム。

【請求項 5】

コンピュータが、
形式の異なる複数の入力ファイルを読み込み、
各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成し、
前記突合キーに基づいて、各中間ファイル内のデータをソートし、
データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成し、

生成した複数の出力ファイルを、データを突合する突合プログラムに入力する
を実行することを特徴とする実行制御方法。

【請求項 6】

形式の異なる複数の入力ファイルを読み込む読込部と、
前記読込部によって読み込まれた各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成する第 1 生成部と、

前記突合キーに基づいて、各中間ファイル内のデータをソートするソート部と、
前記ソート部によってデータがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成する第 2 生成部と、

前記第 2 生成部によって生成された複数の出力ファイルを、データを突合する突合プログラムに入力する入力部と

を有することを特徴とする情報処理装置。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、実行制御プログラム、実行制御方法および情報処理装置に関する。

【背景技術】

【0002】

クラウドコンピューティングの普及に伴い、クラウド上に保存される大量のデータを複数のサーバで分散して処理を実行する分散処理システムが利用されている。分散処理システムとしては、HDFS (Hadoop Distributed File System) と MapReduce とを基盤技術とする Hadoop (登録商標) が知られている。

【0003】

HDFS は、複数のサーバにデータを分散格納するファイルシステムである。MapReduce は、HDFS 上のデータをタスクと呼ばれる単位で分散処理する仕組みであり、Map 処理、Shuffle ソート処理、Reduces 処理を実行する。なお、Map 処理と Reduces 処理は、一般的に Java (登録商標) で開発され、Shuffle ソート処理は、Hadoop に標準装備されている。このような Hadoop では、1 種類の入力を MapReduce で処理して 1 種類の出力を得るのが一般的である。

【0004】

近年では、Hadoop を活用して、Hadoop には標準で装備されない、バッチ処理などの業務

10

20

30

40

50

プログラム（以下、適宜「外部プログラム」と記載する）を効率的に実行することが行われている。外部プログラムは、形式の異なる複数の入力を処理対象とし、一般的にJava（登録商標）以外のプログラムで開発される。

【0005】

例えば、外部プログラムをHadoopで実行する技術として、Hadoopの標準ツールであるHadoop Streamingが知られている。Hadoop Streamingは、Map処理またはReduce処理で外部プログラムを呼び出す技術である。具体的には、Map処理またはReduce処理において、タスク1つにつき外部プログラムを1回呼び出し、外部プログラムの標準出力に処理結果を出力する。

【0006】

また、複数種類の入力を処理するHadoop関連の技術として、reduce side joinが知られている。例えば、突合せ処理の外部プログラムを実行する場合に、入力種類ごとに入力ファイル名、入力フォーマットを処理するクラス、Map処理を実行するクラスを定義してMap処理を実行し、突合せキーと突合せ対象のデータとを対応付けたデータを出力する。続いて、Shuffleソート処理において、突合せキーでソートされ、突合せキーごとにグループ化されたデータを出力する。その後、入力種類ごとにReduceクラスを定義してReduce処理を実行し、マッチング処理を実行する。

【先行技術文献】

【非特許文献】

【0007】

【非特許文献1】Deans and S. Ghemawat, MapReduce著「Simplified Data Processing on Large Clusters, Proceedings of the 6th Symposium on Operating Systems Design and Implementation」、pp.137-150、December 6 2004

【非特許文献2】Apache Hadoop 1.0.3 documentation, 「Hadoop Streaming」, URL「<http://hadoop.apache.org/docs/r1.0.3/streaming.html>」

【非特許文献3】Tom White著、「Hadoop第二版、オライリージャパン、8.3.2 reduce側結合」、P269-272、2011年7月発行

【発明の概要】

【発明が解決しようとする課題】

【0008】

しかしながら、複数入力を処理対象とする外部プログラムをHadoopなどの分散処理システムで実行させるには制約が多いことから、現実的には実行させるのが難しいという問題がある。

【0009】

具体的には、Hadoop Streamingは、標準入出力経由で処理対象のデータを出力するので、引数や環境変数で処理対象のデータを受信する外部プログラムを呼び出して実行することができない。

【0010】

また、reduce side joinでは、既存の外部プログラムと同様の処理を実行するプログラムを再開発してHadoopに移植することになり、再開発のリスクや移植によるリスクがあり、頻繁に外部プログラムを移植することができず、開発性が低い。

【0011】

例えば、reduce side joinを用いる場合には、外部プログラムが処理対象とする入力ごとに、Map処理で処理対象となるMap処理クラスを実装し、さらに、入力ごとにReduce処理で処理対象となるReduce処理クラスを実装する。また、これらのクラスを実装する際に、外部プログラムが処理対象とするRDB (Relational Database)のファイルまたはRDBのファイルをアンロードしたCSV (Comma Separated Values) ファイルとは異なるKVS (Key Value Store) 方式でデータを定義する。また、Java以外で開発されたプログラムであっても、reduce side joinのReduce処理で呼び出せるように、Javaに再開発して移植する。

10

20

30

40

50

【 0 0 1 2 】

つまり、reduce side joinを用いる場合には、複数種類の入力を1種類の入力のように見せかけて、Map処理およびReduce処理を実行し、Reduce処理において移植した外部プログラムを呼び出して実行することになる。

【 0 0 1 3 】

ところが、MapクラスやReduceクラスは、外部プログラムが処理する何百カラムの複雑なデータに対して、装置の支援無しに人手によって実装することになり、時間もかかり、人為的なミスリスクも増える。また、外部プログラムの移植は、リスクが高く、好ましい手法とは言えない。なぜなら、外部プログラムは、複雑な業務ロジックが実装され、既に多くのテストを繰り返して実行実績も豊富なこともあり、簡単に移植できるものではないからである。

10

【 0 0 1 4 】

このように、Hadoop Streamingやreduce side joinを用いて外部プログラムを実行する場合には、作業時間の増加、人為的ミスの増加、外部プログラムの移植に伴うリスクなどがあり、外部プログラムをHadoopで実行するのが難しく、Hadoopの開発性の低下にもなる。

【 0 0 1 5 】

1つの側面では、複数入力を処理対象とする外部プログラムを分散処理システムで実行させる際の制約を緩和することができる実行制御プログラム、実行制御方法および情報処理装置を提供することを目的とする。

20

【課題を解決するための手段】

【 0 0 1 6 】

第1の案では、コンピュータは、形式の異なる複数の入力ファイルを読み込み、各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、入力ファイルごとに生成する。コンピュータは、突合キーに基づいて、各中間ファイル内のデータをソートする。コンピュータは、データがソートされた各中間ファイルから、入力ファイルの各形式のデータをそれぞれ抽出して、各入力ファイルに対してデータがソートされた複数の出力ファイルを生成する。コンピュータは、生成した複数の出力ファイルを、データを突合する突合プログラムに入力する。

【発明の効果】

30

【 0 0 1 7 】

本発明の1実施態様によれば、複数入力を処理対象とする外部プログラムを分散処理システムで実行させる際の制約を緩和することができる。

【図面の簡単な説明】

【 0 0 1 8 】

【図1】図1は、実施例1に係る分散処理システムの全体構成例を示す図である。

【図2】図2は、実施例1に係るマスタ計算機の機能構成を示す機能ブロック図である。

【図3】図3は、設定ファイルDBに記憶される情報の例を示す図である。

【図4】図4は、タスクリストDBに記憶される情報の例を示す図である。

【図5】図5は、スレーブ計算機に送信されるReduceタスク情報の例を示す図である。

40

【図6】図6は、実施例1に係るスレーブ計算機の機能構成を示す機能ブロック図である。

【図7】図7は、実施例1に係るシステムが実行する外部プログラム実行処理の流れを示すシーケンス図である。

【図8】図8は、実施例1に係るMap処理の流れを示すフローチャートである。

【図9】図9は、実施例1に係るReduce処理の流れを示すフローチャートである。

【図10】図10は、実施例2に係るReduce処理の初期化の流れを示すフローチャートである。

【図11】図11は、実施例2に係るReduce処理の本処理および完了処理の流れを示すフローチャートである。

50

【図 1 2】図 1 2 は、実施例 2 に係る出力読み込みスレッドの起動処理の流れを示すフローチャートである。

【図 1 3】図 1 3 は、Map処理の具体例を説明する図である。

【図 1 4】図 1 4 は、シャッフルソート処理の具体例を説明する図である。

【図 1 5】図 1 5 は、Reduce処理の具体例を説明する図である。

【図 1 6】図 1 6 は、突合プログラム実行後の出力先の具体例を説明する図である。

【図 1 7】図 1 7 は、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】

【0019】

以下に、本願の開示する実行制御プログラム、実行制御方法および情報処理装置の実施例を図面に基づいて詳細に説明する。なお、この実施例によりこの発明が限定されるものではない。

【実施例 1】

【0020】

[全体構成]

図 1 は、実施例 1 に係る分散処理システムの全体構成例を示す図である。図 1 に示すように、この分散処理システムは、分析者端末 2、入力 D B (DataBase) サーバ 3、マスタ計算機 10、スレーブ計算機 20、スレーブ計算機 30 がネットワーク 1 を介して相互に通信可能に接続される。

【0021】

この分散処理システムでは、Hadoop (登録商標) などの分散処理フレームワークを使用した分散処理アプリケーションが各計算機で実行されており、データ基盤としてHDFSなどを使用する。分析者端末 2 は、分散処理システムを利用して、分散処理フレームワークとは異なるフレームワークで実装された外部プログラムを実行し、データの解析を行うユーザの端末である。

【0022】

入力 D B サーバ 3 は、外部プログラムが処理対象とする複数種類の入力ファイルを記憶するデータベースサーバである。例えば、入力 D B サーバ 3 は、商品マスタ D B と売上明細 D B とを記憶する。また、入力 D B サーバ 3 が記憶する D B は、R D B のファイルまたは R D B のファイルをアンロードした C S V ファイル等で構成されているものとする。

【0023】

マスタ計算機 10 は、分散処理システムを統括的に管理するサーバである。例えば、マスタ計算機 10 は、どのデータがいずれのスレーブ計算機に格納されているのかを管理し、各スレーブ計算機に割当てするタスクやジョブなどを管理する。

【0024】

スレーブ計算機 20 とスレーブ計算機 30 は、分散処理アプリケーションを実装し、Map処理やReduce処理を実行して、HDFSで管理されるデータを分散処理するサーバである。この各スレーブ計算機は、形式の異なる複数の入力ファイルを読み込み、各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成する。そして、各スレーブ計算機は、突合キーに基づいて、各中間ファイル内のデータをソートする。続いて、各スレーブ計算機は、データがソートされた各中間ファイルから、入力ファイルの各形式のデータをそれぞれ抽出して、各入力ファイルに対してデータがソートされた複数の出力ファイルを生成する。その後、各スレーブ計算機は、生成した複数の出力ファイルを、データを突合する突合プログラムに inputsする。

【0025】

このように、各スレーブ計算機は、Hadoopとはフレームワークが異なる突合プログラムなどの外部プログラムをHadoopで呼び出す際に、突合プログラムの設定に従って各入力ファイルから突合キーを抽出および突合キーでソートして新ファイルを生成する。そして、各スレーブ計算機は、新ファイルを入力ファイルに指定して突合プログラムを呼び出す。

10

20

30

40

50

【 0 0 2 6 】

つまり、各スレーブ計算機は、Hadoopにおいて複数種類の入力ファイルを読み込む際に、外部プログラムに直接読み込ませるという制約を設ける。このことにより、各スレーブ計算機は、外部プログラムの設定を与えることで、複数種類の入力ファイルを実行する外部プログラムを呼び出すことができる。したがって、複数入力を処理対象とする外部プログラムを分散処理システムで実行させる際の制約を緩和することができる。

【 0 0 2 7 】

〔 マスタ計算機の構成 〕

図 2 は、実施例 1 に係るマスタ計算機の機能構成を示す機能ブロック図である。図 2 に示すように、マスタ計算機 1 0 は、通信制御部 1 1、記憶部 1 2、制御部 1 3 を有する。10
なお、通信制御部 1 1 は、例えばネットワークインタフェースカードなどである。記憶部 1 2 は、メモリやハードディスクなどの記憶装置である。制御部 1 3 は、プロセッサなどの電子回路である。

【 0 0 2 8 】

通信制御部 1 1 は、各スレーブ計算機、分析者端末 2、入力 D B サーバ 3 との間で情報を送受信する処理部である。例えば、通信制御部 1 1 は、各スレーブ計算機から Map 要求や Reduce 要求を受信し、各スレーブ計算機に Map タスク情報や Reduce タスク情報を送信する。また、通信制御部 1 1 は、入力 D B サーバ 3 から入力ファイル等を受信する。

【 0 0 2 9 】

記憶部 1 2 は、設定ファイル D B 1 2 a とタスクリスト D B 1 2 b とを記憶する記憶部20
である。また、記憶部 1 2 は、各種処理の途中経過を記憶する一時領域や、分散処理アプリケーション等を記憶する。

【 0 0 3 0 】

設定ファイル D B 1 2 a は、Java フレームワークなどの分散処理フレームワーク以外のフレームワークを用いて実装される外部プログラムの処理内容にしたがって設定された設定ファイルを記憶する。ここで記憶される条件は、管理者等によって設定される。また、外部プログラムは、例えば NetCOBOL (登録商標) などを実装された突合プログラムが該当する。

【 0 0 3 1 】

図 3 は、設定ファイル D B に記憶される情報の例を示す図である。なお、図 3 では、入力30
ファイルが 2 つであり、出力ファイルが 2 つである場合を例示している。入力ファイルとは、突合プログラムが処理対象とする種類の異なる複数のファイルであり、出力ファイルとは、各スレーブ計算機が外部プログラムを実行した結果を格納する H D F S 上の共有ディレクトリ上の共有ファイル である。

【 0 0 3 2 】

図 3 に示すように、設定ファイル D B 1 2 a は、ジョブ設定を行うために設定された各種条件を記憶する。具体的には、設定ファイルには、入力ファイルごとに、入力データファイル名、入力フォーマット処理クラス名、入力フォーマット処理オプション、入力ファイル環境変数名が設定される。

【 0 0 3 3 】

例えば、入力データファイル名 0 1 として「/usr/transaction.txt」が設定されており、これは、入力ファイルの 1 つが「/usr/transaction.txt」であることを示す。また、入力フォーマット処理クラス名 0 1 として「固定長フォーマット」が設定されており、これは、入力ファイルを再構築する際に固定長フォーマットを用いることを示す。また、入力フォーマット処理オプション 0 1 として「レコード長：80BYTE」が設定されており、これは、入力ファイルを再構成する際に 80 バイトの大きさを再構築することを示す。また、入力ファイル環境変数名 0 1 として「TRAN」が設定されており、これは、入力ファイル 0 1 がトランザクションファイルであることを示す。

【 0 0 3 4 】

さらに、設定ファイルには、出力データディレクト名、出力ファイル名、出力ファイ50

ル環境変数名が設定される。例えば、出力データディレクトリ名として「/output」が設定されており、これは、各スレーブ計算機が外部プログラムの処理結果を出力する先のHDFSディレクトリとして「/output」が設定されていることを示す。また、出力ファイル名01として「processed.txt」が設定されており、これは、スレーブ計算機20が正常に終了した処理結果を「/output」のファイル「processed.txt」に出力することを示す。なお、図3には、エラーが発生した場合には、「/output」のファイル「error.txt」に出力することも設定されている。また、出力ファイル環境変数名として「PROCOUT」が設定されており、これは、出力ファイルの環境変数として「PROCOUT」を指定することを示す。

【0035】

さらに、設定ファイルには、Reduceアプリケーションやソートキーが設定される。例えば、Reduceアプリケーションには「a.out %in01 %in02 %out01 %out02」が設定されており、これは、外部プログラムに入力ファイルを渡す際に、1番目の入力ファイルは第1引数で渡し、2番目の入力ファイルは第2引数として渡すことを示す。また、ソートキー01には「4BYTE~8BYTE」が設定されており、これは、入力ファイル01の「4BYTE~8BYTE」をソートキーとして使用することを示す。

【0036】

タスクリストDB12bは、Mapタスク管理部15等によって生成されたMap処理またはReduce処理のタスクリストを記憶する。図4は、タスクリストDBに記憶される情報の例を示す図である。図4に示すように、タスクリストDB12bは、「タスクID、入力情報、クラス名」を記憶する。ここで、「タスクID」は、タスクを識別する識別子であり、「入力情報」は、入力データとして処理する情報を示し、「クラス名」は、Map処理の際に処理対象とするクラス名を示す。

【0037】

例えば、タスクID「task_m_1」には、入力情報としてファイル名「/usr/transaction.txt」、「開始位置=0」、「サイズ=1000」が設定されており、クラス名として「フォーマット処理=本フォーマット処理01」、「マップ処理=本Map処理01」が設定されている。これは、入力フォーマット01を用いたMap処理01では、「/usr/transaction.txt」の0バイト目から1000バイト目までを入力データとして使用することを示している。なお、入力フォーマット01を使用するとは、例えば、図3に示した入力ファイルに関する設定条件のうち「01」に該当する条件を使用することを示す。

【0038】

図2に戻り、制御部13は、設定読込部14、Mapタスク管理部15、Reduceタスク管理部16を有し、これらによって、各スレーブ計算機が実行するMap処理およびReduce処理を管理する処理部である。

【0039】

設定読込部14は、設定ファイルDB12aから設定ファイルを読込んで、Mapタスク管理部15やReduceタスク管理部16に、設定ファイルに設定される各種条件を通知する処理部である。また、設定読込部14は、各スレーブ計算機に対しても、設定ファイルに設定される各種条件を通知する。

【0040】

Mapタスク管理部15は、各スレーブ計算機で実行されるMapタスクを管理する処理部である。例えば、Mapタスク管理部15は、Mapタスク準備として、入力ファイルを先頭から所定サイズごとに区切り、入力ファイルの番号でフォーマット処理およびマップ処理のクラス名を切替えて、タスクリストDB12bに登録する。

【0041】

また、Mapタスク管理部15は、各スレーブ計算機からのタスク要求を受信すると、Mapタスク情報を応答して、各スレーブ計算機にMapタスクを割当てる。例えば、Mapタスク管理部15は、タスクIDが「task_m_1」のレコードをタスクリストDB12bから読み出して、スレーブ計算機20に送信する。

【0042】

10

20

30

40

50

Reduceタスク管理部 16 は、各スレーブ計算機で実行されるReduceタスクを管理する処理部である。具体的には、Reduceタスク管理部 16 は、各スレーブ計算機からのタスク要求を受信すると、Reduceタスク情報を応答して、各スレーブ計算機にReduceタスクを割当てる。図 5 は、スレーブ計算機に送信されるReduceタスク情報の例を示す図である。図 5 に示すように、スレーブ計算機に送信されるReduceタスク情報は、「タスク ID、Map 結果情報、クラス名」を対応付けて記憶する。

【 0 0 4 3 】

「タスク ID」は、Reduceタスクを識別する識別子である。「Map 結果情報」は、Map 処理の結果を特定する情報である。「クラス名」は、Reduce 処理の際に処理対象とするクラス名を示す。図 5 の場合、タスク ID 「task__r__1」のReduceタスクは、スレーブ計算機 20 で処理されたMapタスク ID 「task__m__1」のMap処理の結果とスレーブ計算機 30 で処理されたMapタスク ID 「task__m__2」のMap処理の結果とを用いて、Reduce処理を実行することを示す。なお、タスク ID 「task__r__1」のReduceタスクは、クラス名として「本Reduceクラス」が設定されている。

【 0 0 4 4 】

[スレーブ計算機の構成]

次に、スレーブ計算機について説明するが、スレーブ計算機 20 とスレーブ計算機 30 とは同様の構成を有するので、ここでは、スレーブ計算機 20 について説明する。図 6 は、実施例 1 に係るスレーブ計算機の機能構成を示す機能ブロック図である。図 6 に示すように、スレーブ計算機 20 は、通信制御部 21、記憶部 22、制御部 23 を有する。なお、通信制御部 21 は、例えばネットワークインタフェースカードなどである。記憶部 22 は、メモリやハードディスクなどの記憶装置である。制御部 23 は、プロセッサなどの電子回路である。

【 0 0 4 5 】

通信制御部 21 は、マスタ計算機 10、入力 DB サーバ 3、分析者端末 2 との間で情報を送受信する処理部である。例えば、通信制御部 21 は、マスタ計算機 10 のMapタスク要求やReduceタスク要求を送信し、マスタ計算機 10 からMapタスク情報やReduceタスク情報を受信する。また、通信制御部 21 は、入力 DB サーバ 3 から入力ファイル等を受信する。

【 0 0 4 6 】

記憶部 22 は、中間ファイル DB 22 a と一時ファイル DB 22 b とを記憶する記憶部である。また、記憶部 22 は、分散処理アプリケーションや外部プログラム、入力 DB サーバ 3 から取得された入力ファイルや入力データ等を記憶する。

【 0 0 4 7 】

中間ファイル DB 22 a は、制御部 23 等によって読み込まれた入力ファイルや入力データ、また、Map処理の結果やReduce処理の結果を記憶する。すなわち、中間ファイル DB 22 a は、スレーブ計算機 20 が外部プログラムを呼び出す前に実行した処理の結果等を記憶する。一時ファイル DB 22 b は、外部プログラムによって処理された処理結果を記憶する。

【 0 0 4 8 】

制御部 23 は、Map処理部 24 とReduce処理部 25 を有し、これらによって、Map処理およびReduce処理を実行し、さらに、外部プログラムを呼び出して実行する処理部である。Map処理部 24 は、Map初期化部 24 a、フォーマット初期化部 24 b、抽出部 24 c を有し、これらによってMap処理を実行する処理部である。

【 0 0 4 9 】

Map初期化部 24 a は、マスタ計算機 10 から受信したMapタスク情報にしたがって、Mapタスクの初期化を実行する処理部である。具体的には、Map初期化部 24 a は、どのようなフォーマットで入力ファイルを読み込み、また、ソートキーをどのように読み込むかを、設定ファイルにしたがって設定する。

【 0 0 5 0 】

例えば、Map初期化部 2 4 a は、受信したMapタスク情報から「本Map処理 0 1」を抽出する。続いて、Map初期化部 2 4 a は、抽出した「本Map処理 0 1」から「0 1」を抽出し、「0 1」に対応する「ソートキー 0 1 = 4 BYTE ~ 8 BYTE」を設定ファイルから取得する。そして、Map処理部 2 4 は、取得した情報をフォーマット初期化部 2 4 b、抽出部 2 4 c、Reduce処理部 2 5 等に出力する。

【 0 0 5 1 】

フォーマット初期化部 2 4 b は、Map処理実行結果を出力する中間ファイルのフォーマットの初期化を実行する処理部である。具体的には、フォーマット初期化部 2 4 b は、複数の入力を読み込むためのアダプタを入力ファイルごとに割当てて。

【 0 0 5 2 】

例えば、フォーマット初期化部 2 4 b は、Map初期化部 2 4 a から受信した「本Map処理 0 1」の「0 1」に対応する入力フォーマット処理クラス名や入力フォーマット処理オプション等を設定ファイルから取得して初期化を実行する。ここでは、フォーマット初期化部 2 4 b は、「入力フォーマット処理クラス名 0 1 = 固定長フォーマット」と、「入力フォーマット処理オプション 0 1 = レコード長 (80BYTE)」等を取得する。

【 0 0 5 3 】

抽出部 2 4 c は、Map処理を実行し、入力ファイル等から該当するデータを抽出して中間ファイルを生成する処理部である。具体的には、抽出部 2 4 c は、フォーマット初期化部 2 4 b によって設定されたフォーマットで入力ファイルの各レコードを読み込む。そして、抽出部 2 4 c は、入力からソートキーを抽出してMap出力の K e y とする。また、抽出部 2 4 c は、何番目の入力であったかを示すファイルインデックスを K e y または V a l u e の一部とし、読み込まれたレコード全体を K e y または V a l u e の一部として、Map処理結果として出力する。

【 0 0 5 4 】

例えば、抽出部 2 4 c は、Map初期化部 2 4 a から受信した「ソートキー 0 1 = 4 BYTE ~ 8 BYTE」にしたがって、入力ファイルのレコード (入力行) からソートキーを抽出する。続いて、抽出部 2 4 c は、抽出したソートキーと、ソートキーの抽出元の入力行とを対応付けた中間ファイルを生成して、中間ファイル D B 2 2 a に格納する。つまり、抽出部 2 4 c は、ソートキーを K e y、入力行を V a l u e とする K V S 形式の中間ファイルを生成する。なお、マスタ計算機 1 0 は、ここで格納された中間ファイルを読み出して、Map処理結果を管理する。

【 0 0 5 5 】

Reduce処理部 2 5 は、Shuffle処理部 2 5 a、Reduce初期化部 2 5 b、再構築部 2 5 c を有し、これらによって外部プログラムを呼び出してReduce処理を実行する処理部である。すなわち、Reduce処理部 2 5 は、外部プログラムに適した複数の出力ファイルを生成して、外部プログラムを呼び出す。

【 0 0 5 6 】

Shuffle処理部 2 5 a は、Hadoopのreduce side joinなどで実行されるシャッフルソート処理を実行する処理部である。例えば、Shuffle処理部 2 5 a は、マスタ計算機 1 0 に対してReduceタスク要求を送信して、Reduceタスク情報を受信する。そして、Shuffle処理部 2 5 a は、受信したReduceタスク情報のMap結果情報にしたがって、自装置がシャッフルソート処理の対象とするMap処理結果を各スレーブ計算機から収集して中間ファイル D B 2 2 a に格納する。一例を挙げると、Shuffle処理部 2 5 a は、Mapタスク I D 「task_m_2」のMap処理の結果をスレーブ計算機 3 0 から取得し、Mapタスク I D 「task_m_1」のMap処理の結果を自装置内の中間ファイル D B 2 2 a から取得する。なお、割当てられる手法は、Hadoopの分散手法を用いる。

【 0 0 5 7 】

その後、Shuffle処理部 2 5 a は、抽出部 2 4 c によって抽出されたソートキーを用いて、収集したMap処理の結果をソートする。そして、Shuffle処理部 2 5 a は、ソートした結果を同一キーでグループ化して、中間ファイル D B 2 2 a に格納する。なお、各スレー

10

20

30

40

50

ブ計算機で分散処理されたシャッフルソート処理結果は、マスタ計算機 10 によって各スレーブ計算機から収集されて管理される。

【 0 0 5 8 】

Reduce初期化部 2 5 b は、Shuffle処理部 2 5 a が取得したReduceタスク情報にしたがって、Reduce処理を実行する前段階として初期化を実行する処理部である。具体的には、Reduce初期化部 2 5 b は、どのように外部プログラムを呼び出し、入力ファイルをどのように渡すかを設定する。

【 0 0 5 9 】

例えば、Reduce初期化部 2 5 b は、「a.out %in01 %in02 %out01 %out02」と指定して、1 番目の入力ファイルを外部プログラムの第 1 引数に渡す。あるいは、Reduce初期化部 2 5 b は、「ENVNAME.01=TRAN」と指定して、入力ファイル 0 1 のファイル名をTRAN環境変数に設定して、外部プログラムを呼び出すと設定する。

【 0 0 6 0 】

また、Reduce初期化部 2 5 b は、外部プログラムへの入力に該当するファイル名を決定し、外部プログラムを決定したファイル名を用いて呼び出すための準備を実行する。例えば、Reduce初期化部 2 5 b は、呼出コマンドの引数の「%in01」を「./in01」に置き換えたり、ユーザ設定の環境変数TRANにファイル名「./in01」を設定したりする。

【 0 0 6 1 】

再構築部 2 5 c は、Reduce初期化部 2 5 b によって初期化された情報を用いて、マスタ計算機 10 から指定されたReduceタスクを実行して出力ファイルを生成する。そして、再構築部 2 5 c は、出力ファイルを入力ファイルに指定して、外部プログラムを実行する処理部である。すなわち、Reduce初期化部 2 5 b は、マスタ計算機 10 から指定されたReduceタスクを実行し、Shuffle処理部 2 5 a の処理結果から、該当するレコードを読み出して該当するファイルに書き込む。

【 0 0 6 2 】

このとき、再構築部 2 5 c は、Shuffle処理部 2 5 a から出力された K V S 形式の処理結果を、元の入力ファイルと同じファイル形式に再変換する。そして、再構築部 2 5 c は、Reduce初期化部 2 5 b によって設定された当該ファイルの位置等を引数や環境変数を指定して、外部プログラムを呼び出して実行する。

【 0 0 6 3 】

また、再構築部 2 5 c は、外部プログラムの実行によって得られた結果を、一時ファイル 2 2 b に格納する。このとき、再構築部 2 5 c は、設定ファイル等にしたがって、正常結果とエラー結果とを区別して、出力することもできる。

【 0 0 6 4 】

[シーケンス]

次に、図 7 を用いて、図 1 に示した分散処理システムで外部プログラムを実行する際の全体的な処理の流れを説明する。図 7 は、実施例 1 に係るシステムが実行する外部プログラム実行処理の流れを示すシーケンス図である。なお、ここで説明を簡略化するために、1 台のスレーブ計算機 20 を例にして説明する。

【 0 0 6 5 】

図 7 に示すように、マスタ計算機 10 の設定読込部 1 4 は、設定ファイル D B 1 2 a から設定情報を読込む (S 1 0 1)。続いて、Mapタスク管理部 1 5 は、入力 D B サーバ 3 から入力ファイルが予め読み込まれた H D F S などの共有ファイルシステムから入力ファイルを 1 つ読み込み (S 1 0 2)、データを所定サイズごとに分割し (S 1 0 3)、タスクリストに登録する (S 1 0 4)。例えば、Mapタスク管理部 1 5 は、入力ファイルの番号で、フォーマット処理およびMap処理のクラス名を切替えて登録する。

【 0 0 6 6 】

そして、マスタ計算機 10 のMapタスク管理部 1 5 は、入力ファイルが他にも存在する場合には (S 1 0 5 : Y e s)、S 1 0 2 に戻って以降の処理を繰り返す。一方、Mapタスク管理部 1 5 は、入力ファイルが他に存在しない場合には (S 1 0 5 : N o)、事前準

10

20

30

40

50

備が終了したことを示す通知を各スレーブ計算機に送信する（S 1 0 6 と S 1 0 7）。なお、スレーブ計算機が定期的にタスク要求 S 1 0 8 を実行するように構成したり、スレーブ計算機がタスク要求 S 1 0 8 を実行した場合に、S 1 1 0 の Map タスク応答が返らないように構成することもでき、このような場合は、S 1 0 6 および S 1 0 7 を省略することもできる。

【 0 0 6 7 】

この終了通知を受信したスレーブ計算機 2 0 の Map 処理部 2 4 は、Map タスク要求をマスタ計算機 1 0 に送信する（S 1 0 8 と S 1 0 9）。この要求を受信したマスタ計算機 1 0 の Map タスク管理部 1 5 は、図 4 のタスクリストから該当するタスクを抽出して、Map タスクとしてスレーブ計算機 2 0 に応答する（S 1 1 0 と S 1 1 1）。

10

【 0 0 6 8 】

スレーブ計算機 2 0 の Map 処理部 2 4 は、受信した Map タスクにしたがって Map 処理を実行する（S 1 1 2）。そして、Reduce 処理部 2 5 は、Map 処理が終了すると、Reduce タスク要求をマスタ計算機 1 0 に送信する（S 1 1 3 と S 1 1 4）。この要求を受信したマスタ計算機 1 0 の Reduce タスク管理部 1 6 は、図 5 に示す Reduce タスクをスレーブ計算機 2 0 に応答する（S 1 1 5 と S 1 1 6）。その後、スレーブ計算機 2 0 の Reduce 処理部 2 5 は、受信した Reduce タスクにしたがって Reduce 処理を実行する（S 1 1 7）。

【 0 0 6 9 】

（Map 処理）

次に、図 7 に示した Map 処理について説明する。図 8 は、実施例 1 に係る Map 処理の流れを示すフローチャートである。図 8 に示すように、スレーブ計算機 2 0 の Map 初期化部 2 4 a は、クラス名、または、クラス特有の定義より、入力番号（N N）を抽出する（S 2 0 1）。例えば、Map 初期化部 2 4 a は、マスタ計算機 1 0 から受信した Map タスク情報から「本 Maps 処理 0 1」の「0 1」を上記 N N として抽出する。

20

【 0 0 7 0 】

続いて、Map 初期化部 2 4 a は、入力番号（N N）とジョブ設定（Map タスク情報）とからソートキーの読み込み位置を抽出する（S 2 0 2）。例えば、Map 初期化部 2 4 a は、N N = 0 1 であることから、ソートキー（N N）= ソートキー（0 1）と特定し、「ソートキー = 0 1」に対応付けられる「4BYTE ~ 8BYTE」を抽出する。

【 0 0 7 1 】

30

その後、Map 処理部 2 4 のフォーマット初期化部 2 4 b は、クラス名、または、クラス特有の定義より、入力番号（N N）を抽出する（S 2 0 3）。例えば、フォーマット初期化部 2 4 b は、S 2 0 1 と同様の手法で、「0 1」を上記 N N として抽出する。

【 0 0 7 2 】

続いて、フォーマット初期化部 2 4 b は、入力番号（N N）とジョブ設定（Map タスク情報）とから入力フォーマットのオプションを抽出し初期化を実行する（S 2 0 4）。例えば、フォーマット初期化部 2 4 b は N N = 0 1 であることから、「入力フォーマット処理オプション 0 1」に対応付けられる「レコード長：80BYTE」を抽出し、入力ファイルのレコード長を初期化する。

【 0 0 7 3 】

40

その後、抽出部 2 4 c は、入力ファイルから行であるレコードを読み込み、読み込んだ入力行からソートキーの値を抽出する（S 2 0 5）。そして、抽出部 2 4 c は、ソートキー（K e y）を「ソートキーの値」、値（V a l u e）を「N N、入力行」とする K V S 形式の中間ファイルを生成して、中間ファイル D B 2 2 a に出力する（S 2 0 6）。

【 0 0 7 4 】

（Reduce 処理）

次に、図 7 に示した Reduce 処理について説明する。図 9 は、実施例 1 に係る Reduce 処理の流れを示すフローチャートである。図 9 に示すように、Reduce 処理部 2 5 の Shuffle 処理部 2 5 a は、各スレーブ計算機から Map タスク結果を収集し（S 3 0 1）、収集した結果をソートし（S 3 0 2）、ソートした結果を同一キーでグループ化する（S 3 0 3）。

50

【 0 0 7 5 】

続いて、Reduce初期化部 2 5 b は、全ての入力番号 (N N) について、 S 3 0 4 から S 3 1 1 を実行する。具体的には、Reduce初期化部 2 5 b は、入力 N N 用一時ファイルを一時領域である一時ファイル D B 2 2 b にオープン (生成) する (S 3 0 5)。続いて、Reduce初期化部 2 5 b は、フォーマット処理 N N に対応する書き込みフォーマット処理を入力 N N 用一時ファイルと関連付ける (S 3 0 6)。

【 0 0 7 6 】

その後、Reduce初期化部 2 5 b は、書き込みフォーマット処理を入力フォーマット処理 N N のオプションで初期化する (S 3 0 7)。例えば、Reduce初期化部 2 5 b は、入力フォーマット処理オプション 0 1 に対応する「レコード長 : 80BYTE」に、出力先のファイルを初期化する。

10

【 0 0 7 7 】

続いて、Reduce初期化部 2 5 b は、入力ファイル番号 N N、一時ファイル名、出力フォーマットクラスを対応付けた入力対応表を記憶部 2 2 等に登録する (S 3 0 8)。例えば、Reduce初期化部 2 5 b は、「入力ファイル番号、ファイル名、出力フォーマットクラス」として「 0 1、./tmp/in01.txt、固定長フォーマット出力クラス」や「 0 2、./tmp/in02.txt、改行ありフォーマット出力クラス」を生成する。

【 0 0 7 8 】

その後、Reduce初期化部 2 5 b は、ジョブ設定 (Reduceタスク情報) で指定された Reduceアプリケーションの引数の「%inNN」を入力 N N 用一時ファイルのファイル名で置き換える (S 3 0 9)。例えば、Reduce初期化部 2 5 b は、Reduceアプリケーション「a.out %in01 %in02」を「a.out ./tmp/in01.txt ./tmp/in02.txt」に書き換える。

20

【 0 0 7 9 】

そして、Reduce初期化部 2 5 b は、ジョブ設定で指定された入力ファイル環境変数名 N N を参照して、環境変数に入力 N N 用一時ファイルのファイル名を設定する (S 3 1 0)。例えば、Reduce初期化部 2 5 b は、環境変数の「TRAN」に「./tmp/in01.txt」、「MASTER」に「./tmp/in02.txt」を設定する。

【 0 0 8 0 】

その後、再構築部 2 5 c は、全てのキーの値のリストで S 3 1 2 ~ S 3 1 6 を実行する。具体的には、再構築部 2 5 c は、値 = { N N、入力行 } として N N (入力ファイル番号) と入力行を抽出する (S 3 1 3)。

30

【 0 0 8 1 】

そして、再構築部 2 5 c は、入力対応表と N N とから出力フォーマットクラスオブジェクトを抽出し (S 3 1 4)、出力フォーマットクラスオブジェクト経由で入力行を一時ファイルに出力する (S 3 1 5)。このとき、実際にファイルの中身が出力されることとなる。

【 0 0 8 2 】

その後、再構築部 2 5 c は、Reduceアプリケーションの文字列と環境変数で、外部プログラムを呼び出して実行し、終了するまで待機する (S 3 1 7)。そして、再構築部 2 5 c は、カレントディレクトリの全てあるいは一部のファイルをジョブ設定の出力ディレクトリのサブディレクトリにコピーする (S 3 1 8)。

40

【 0 0 8 3 】

このようにすることで、外部プログラムをHadoopに移植することもなく、複雑で膨大なタスク定義を実装することもなく、Hadoopで外部プログラムを呼び出して実行することができる。この結果、作業時間の削減、人為的ミスの削減、外部プログラムの移植に伴うリスクが低減でき、Hadoopの開発性を向上させることもできる。

【実施例 2】

【 0 0 8 4 】

実施例 1 では、外部プログラムの入力ファイルをディスク上のファイルとする例で説明したが、これに限定されるものではなく、例えば、いわゆる名前付きパイプ (named pip

50

e)を使用することもできる。また、外部プログラムの出力ファイルを取得し、MapReduceの入力として使用できる共有ファイルシステム上にコピーすることもできる。

【0085】

そこで、実施例2では、名前付きパイプおよび共有ファイルシステムを使用する例を説明する。なお、マスタ計算機10が実行する処理、各スレーブ計算機が実行するMap処理とShuffle処理については、実施例1と同様なので、説明を省略する。ここでは、実施例1とは異なるReduce処理について説明する。なお、図10のS401の前に図9のS301からS303と同様の処理が実行される。

【0086】

(Reduce処理の初期化処理)

10

図10は、実施例2に係るReduce処理の初期化の流れを示すフローチャートである。図10に示すように、スレーブ計算機20のReduce初期化部25bは、全ての入力番号(NN)について、S401からS409を実行する。

【0087】

具体的には、Reduce初期化部25bは、入力用名前付きパイプNNを一時領域にオープンする(S402)。ここで一時領域とは、例えば、スレーブ計算機20の記憶部22の一時ファイルDB22bなどである。

【0088】

続いて、Reduce初期化部25bは、フォーマット処理NNに対応する出力フォーマット処理を入力NN用一時ファイルと関連付ける(S403)。その後、Reduce初期化部25bは、出力フォーマット処理をフォーマット処理NNのオプションで初期化する(S404)。

20

【0089】

続いて、Reduce初期化部25bは、入力用名前付きパイプの書き込み用スレッドを立ち上げる(S405)。その後、Reduce初期化部25bは、入力ファイル番号NN、入力用名前付きパイプNNのファイル名、出力フォーマットクラス、書き込みスレッドIDを対応付けた入力対応表を記憶部22等に登録する(S406)。

【0090】

例えば、Reduce初期化部25bは、「入力ファイル番号、ファイル名、出力フォーマットクラス、書き込みスレッドID」として「01、./tmp/in01.txt、固定長フォーマット出力クラス、スレッド101」を生成する。また、Reduce初期化部25bは、「02、./tmp/in02.txt、改行ありフォーマット出力クラス、スレッド102」を生成する。

30

【0091】

その後、Reduce初期化部25bは、ジョブ設定で指定されたReduceアプリケーションの引数の「%inNN」を入力用名前付きパイプNNのファイル名で置き換える(S407)。例えば、Reduce初期化部25bは、Reduceアプリケーション「a.out %in01 %in02 %out01 %out02」を「a.out ./tmp/in01.txt ./tmp/in02.txt %out01 %out02」に書き換える。

【0092】

そして、Reduce初期化部25bは、ジョブ設定で指定された入力ファイル環境変数名NNを参照して、環境変数に入力用名前付きパイプNNのファイル名を設定する(S408)。例えば、Reduce初期化部25bは、環境変数の「TRAN」に「./tmp/in01.txt」、「MASTER」に「./tmp/in02.txt」を設定する。

40

【0093】

その後、Reduce初期化部25bは、全ての出力番号(NN)について、S410からS419を実行する。具体的には、Reduce初期化部25bは、出力用名前付きパイプNNを一時領域にオープンする(S411)。ここで一時領域とは、例えば、スレーブ計算機20の記憶部22の一時ファイルDB22bなどである。

【0094】

続いて、Reduce初期化部25bは、ジョブ設定の出力フォーマット処理NNに対応する

50

出力フォーマット処理クラスを出力ファイルNNと関連付ける（S 4 1 2）。その後、Reduce初期化部 2 5 bは、出力フォーマット処理を出力フォーマット処理NNのオプションで初期化する（S 4 1 3）。

【0 0 9 5】

続いて、Reduce初期化部 2 5 bは、読み込んだ出力を書き出す共有ファイルのファイル名を、ジョブ設定の出力データディレクトリ名と出力ファイル名NNとから生成する（S 4 1 4）。その後、Reduce初期化部 2 5 bは、出力読み込みスレッドの起動処理を実行する（S 4 1 5）。

【0 0 9 6】

そして、Reduce初期化部 2 5 bは、S 4 1 5の起動処理が終了すると、S 4 1 6を実行する。すなわち、Reduce初期化部 2 5 bは、出力ファイル番号NN、出力用名前付きパイプNNのファイル名、出力フォーマットクラス、出力読み込みスレッドID、共有ファイルのファイル名を対応付けた出力対応表を記憶部 2 2等に登録する。例えば、Reduce初期化部 2 5 bは、「0 1、./tmp/out01.txt、改行ありフォーマット出力クラス、スレッド 1 0 3、/output/processed.txt」や「0 2、./tmp/out02.txt、固定長フォーマット出力クラス、スレッド 1 0 4、/output/error.txt」を生成する。

【0 0 9 7】

その後、Reduce初期化部 2 5 bは、ジョブ設定で指定されたReduceアプリケーションの引数の「%outNN」を出力用名前付きパイプNNのファイル名で置き換える（S 4 1 7）。例えば、Reduce初期化部 2 5 bは、Reduceアプリケーション「a.out ./tmp/in01.txt ./tmp/in02.txt %out01 %out02」を「a.out ./tmp/in01.txt ./tmp/in02.txt ./tmp/out01.txt ./tmp/out02.txt」に書き換える。

【0 0 9 8】

そして、Reduce初期化部 2 5 bは、ジョブ設定で指定された出力ファイル環境変数名NNを参照して、環境変数に出力用名前付きパイプNNのファイル名を設定する（S 4 1 8）。その後、再構築部 2 5 cは、Reduceアプリケーションの文字列と環境変数で、外部プログラムを呼び出す（S 4 2 0）。

【0 0 9 9】

（Reduce処理の本処理および完了処理）

続いて、Reduce処理の本処理および完了処理を説明する。図 1 1は、実施例 2に係るReduce処理の本処理および完了処理の流れを示すフローチャートである。図 1 1に示すように、スレーブ計算機 2 0の再構築部 2 5 cは、全てのキーの値のリストでS 5 0 1～S 5 0 5を実行する。具体的には、再構築部 2 5 cは、値 = { NN、入力行 }としてNN（入力ファイル番号）と入力行を抽出する（S 5 0 2）。

【0 1 0 0】

そして、再構築部 2 5 cは、入力対応表とNNとから出力フォーマットクラスオブジェクトを抽出し（S 5 0 3）、出力フォーマットクラスオブジェクト経由で入力行を名前付きパイプに書き込む（S 5 0 4）。

【0 1 0 1】

その後、再構築部 2 5 cは、入力用名前付きパイプをクローズし（S 5 0 6）、Reduceアプリケーションが終了するまで待機する（S 5 0 7）。そして、再構築部 2 5 cは、入力用名前付きパイプの書き込みスレッド全てを終了する（S 5 0 8）。

【0 1 0 2】

続いて、再構築部 2 5 cは、出力用名前付きパイプの読み込みスレッド全てが終了するまで待機する（S 5 0 9）。その後、再構築部 2 5 cは、カレントディレクトリの全てあるいは一部のファイルをジョブ設定の出力ディレクトリのサブディレクトリにコピーする（S 5 1 0）。

【0 1 0 3】

（スレッド起動処理）

続いて、図 1 0のS 4 1 5に示したスレッド起動処理を説明する。図 1 2は、実施例 2

10

20

30

40

50

に係る出力読み込みスレッドの起動処理の流れを示すフローチャートである。図 1 2 に示すように、Reduce初期化部 2 5 b は、共有ファイルのファイル名を書き込みモードでオープンし (S 6 0 1) 、共有ファイルを出力フォーマット処理クラスと関連付けて、初期化する (S 6 0 2) 。

【 0 1 0 4 】

その後、Reduce初期化部 2 5 b は、名前付きパイプを読み込みモードでオープンし (S 6 0 3) 、名前付きパイプを出力フォーマットに対応する入力フォーマット読み込みクラスと関連付けて、初期化する (S 6 0 4) 。

【 0 1 0 5 】

そして、再構築部 2 5 c は、名前付きパイプからデータが読み込める間、 S 6 0 5 ~ S 6 0 8 を繰り返す。具体的には、再構築部 2 5 c は、名前付きパイプから入力フォーマット読み込みクラスオブジェクト経由で 1 行読み込む (S 6 0 6) 。その後、再構築部 2 5 c は、出力フォーマットクラスオブジェクト経由で入力行を共有ファイルに出力する (S 6 0 7) 。

【実施例 3】

【 0 1 0 6 】

次に、入力ファイル 0 1 として商品マスタファイル、入力ファイル 0 2 として売上明細ファイルを用いて、Hadoopで使用されるJavaフレームワークとは異なるNetCOBOLで実装された突合プログラムをHadoop内で読み出して実行する例を説明する。なお、実施例 3 では、1 台のスレーブ計算機に 1 つのMap処理が割当てられている例で説明するが、これに限

【 0 1 0 7 】

図 1 3 は、Map処理の具体例を説明する図であり、図 1 4 は、シャッフルソート処理の具体例を説明する図であり、図 1 5 は、Reduce処理の具体例を説明する図であり、図 1 6 は、突合プログラム実行後の出力先の具体例を説明する図である。

【 0 1 0 8 】

図 1 3 に示すように、各スレーブ計算機がアクセス可能なHDFSディレクトリには、突合プログラムが処理対象とする入力ファイル 0 1 と入力ファイル 0 2 とが格納されている。

【 0 1 0 9 】

入力ファイル 0 1 は、商品マスタであり、「商品 I D、商品名、単価、限定割引」として「0001、お茶、140、 - 」、「0011、梅おにぎり、110、10%」、「0012、鮭おにぎり、120、 - 」を記憶する。また、入力ファイル 0 2 は、売上明細ファイルであり、「伝票 I D、商品 I D、個数」として「0001、0012、2」、「0001、0001、1」、「0002、0011、1」を記憶する。

【 0 1 1 0 】

また、各スレーブ計算機は、「入力 0 1 : 行順ファイルカラム 1 抽出、入力 0 2 : 行順ファイル 2 抽出」が記述された設定ファイルを読み込む。このような状態において、スレーブ計算機 2 0 は、入力ファイル 0 1 を処理するMapタスクがマスタ計算機 1 0 より指定されている。このスレーブ計算機 2 0 は、設定ファイル「入力 0 1 : 行順ファイルカラム 1 抽出」にしたがって、入力ファイル 0 1 のカラム 1 である「商品 I D」をキーに設定する。

【 0 1 1 1 】

そして、スレーブ計算機 2 0 は、各レコードから商品 I D をキーとして抽出し、抽出したキーと抽出元の入力ファイルを示す情報と抽出元のレコードとを対応付けた K V S 形式の中間ファイルを生成する。例えば、スレーブ計算機 2 0 は、商品 I D が 0001 のレコードについて「Key、Value」として「K (0001) 、V (1、0001、お茶、140、 -) 」を生成して、中間ファイルに格納する。

【 0 1 1 2 】

同様に、スレーブ計算機 3 0 は、入力ファイル 0 2 を処理するMapタスクがマスタ計算

10

20

30

40

50

機 1 0 より指定されている。このスレーブ計算機 3 0 は、設定ファイル「入力 0 2 : 行順ファイルカラム 2 抽出」にしたがって、入力ファイル 0 2 のカラム 2 である「商品 I D」をキーに設定する。

【 0 1 1 3 】

そして、スレーブ計算機 3 0 は、各レコードから商品 I D をキーとして抽出し、抽出したキーと抽出元の入力ファイルを示す情報と抽出元のレコードとを対応付けた K V S 形式の中間ファイルを生成する。例えば、スレーブ計算機 3 0 は、伝票 I D が 0001 のレコードについては、商品 I D 「0012」をキーとして抽出し、「Key、Value」として「K (0012) 、 V (2、0001、0012、2) 」を生成して、中間ファイルに格納する。

【 0 1 1 4 】

その後、図 1 4 に示すように、スレーブ計算機 2 0 は、キーが「0001 - 0010」の範囲のデータを担当することがマスタ計算機 1 0 によって指定されており、該当するデータを取得してシャッフルソートを実行する。具体的には、このスレーブ計算機 2 0 は、自装置の Map 処理結果とスレーブ計算機 3 0 の Map 処理結果から、「K e y」が「0001 - 0010」の範囲内であるレコードを読み出して、当該 K e y でソートする。

【 0 1 1 5 】

例えば、スレーブ計算機 2 0 は、自装置の Map 処理結果から「K (0001) 、 V (1、0001、お茶、140、 -) 」を取得し、スレーブ計算機 3 0 の Map 処理結果から「K (0001) 、 V (2、0001、0001、1) 」を取得する。そして、スレーブ計算機 2 0 は、これらをソートして、K (0001) に V { 入力 (1) 、商品 I D (0001) 、商品名 (お茶) 、単価 (140) 、限定割引 (-) } を対応付けたレコードを中間ファイルに出力する。同様に、スレーブ計算機 2 0 は、K (0001) に V { 入力 (2) 、伝票 I D (0001) 、商品 I D (0001) 、個数 (1) } を対応付けたレコードを中間ファイルに出力する。

【 0 1 1 6 】

同様に、スレーブ計算機 3 0 は、キーが「0011 - 0020」の範囲のデータを担当することがマスタ計算機 1 0 によって指定されており、該当するデータを取得してシャッフルソートを実行する。具体的には、このスレーブ計算機 3 0 は、自装置の Map 処理結果とスレーブ計算機 3 0 の Map 処理結果から、「K e y」が「0011 - 0020」の範囲内であるレコードを読み出して、当該 K e y でソートする。

【 0 1 1 7 】

例えば、スレーブ計算機 3 0 は、スレーブ計算機 2 0 の Map 処理結果から「K (0011) 、 V (1、0011、梅おにぎり、110、10%) 」と「K (0012) 、 V (1、0012、鮭おにぎり、120、 -) 」とを取得する。また、スレーブ計算機 3 0 は、自装置の Map 処理結果から「K (0012) 、 V (2、0001、0012、2) 」と「K (0011) 、 V (2、0002、0011、1) 」とを取得する。

【 0 1 1 8 】

そして、スレーブ計算機 3 0 は、これらをソートして、K (0011) に V { 入力 (1) 、商品 I D (0011) 、商品名 (梅おにぎり) 、単価 (110) 、限定割引 (10%) } を対応付けたレコードを中間ファイルに出力する。同様に、スレーブ計算機 3 0 は、K (0011) に、V { 入力 (2) 、伝票 I D (0002) 、商品 I D (0011) 、個数 (1) } を対応付けたレコードを中間ファイルに出力する。

【 0 1 1 9 】

また、スレーブ計算機 3 0 は、これらをソートして、K (0012) に V { 入力 (1) 、商品 I D (0012) 、商品名 (鮭おにぎり) 、単価 (120) 、限定割引 (-) } を対応付けたレコードを中間ファイルに出力する。同様に、スレーブ計算機 3 0 は、K (0012) に V { 入力 (2) 、伝票 I D (0001) 、商品 I D (0012) 、個数 (2) } を対応付けたレコードを中間ファイルに出力する。

【 0 1 2 0 】

その後、図 1 5 に示すように、スレーブ計算機 2 0 は、シャッフルソートした結果に対して Reduce 処理を実行する。具体的には、スレーブ計算機 2 0 は、「K (0001) 、 V { 入

10

20

30

40

50

力(1)、商品ID(0001)、商品名(お茶)、単価(140)、限定割引(-)}」からV{入力(1)、商品ID(0001)、商品名(お茶)、単価(140)、限定割引(-)}を抽出する。そして、スレーブ計算機20は、抽出したV{入力(1)、商品ID(0001)、商品名(お茶)、単価(140)、限定割引(-)}から入力(1)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機20は、入力ファイル01用の名前付きパイプに、変換後の「商品ID(0001)、商品名(お茶)、単価(140)、限定割引(-)」を入力する。

【0121】

同様に、スレーブ計算機20は、「K(0001)、V{入力(2)、伝票ID(0001)、商品ID(0001)、個数(1)}」からV{入力(2)、伝票ID(0001)、商品ID(0001)、個数(1)}を抽出する。そして、スレーブ計算機20は、抽出したV{入力(2)、伝票ID(0001)、商品ID(0001)、個数(1)}から入力(2)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機20は、入力ファイル02用の名前付きパイプに、変換後の「伝票ID(0001)、商品ID(0001)、個数(1)」を入力する。

【0122】

また、スレーブ計算機30は、「K(0011)、V{入力(1)、商品ID(0011)、商品名(梅おにぎり)、単価(110)、限定割引(10%)}」からV{入力(1)、商品ID(0011)、商品名(梅おにぎり)、単価(110)、限定割引(10%)}を抽出する。そして、スレーブ計算機30は、抽出したV{入力(1)、商品ID(0011)、商品名(梅おにぎり)、単価(110)、限定割引(10%)}から入力(1)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機30は、入力ファイル01用の名前付きパイプに、変換後の「商品ID(0011)、商品名(梅おにぎり)、単価(110)、限定割引(10%)」を入力する。

【0123】

また、スレーブ計算機30は、「K(0011)、V{入力(2)、伝票ID(0002)、商品ID(0011)、個数(1)}」からV{入力(2)、伝票ID(0002)、商品ID(0011)、個数(1)}を抽出する。そして、スレーブ計算機30は、抽出したV{入力(2)、伝票ID(0002)、商品ID(0011)、個数(1)}から入力(2)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機30は、入力ファイル02用の名前付きパイプに、変換後の「伝票ID(0002)、商品ID(0011)、個数(1)」を入力する。

【0124】

また、スレーブ計算機30は、「K(0012)、V{入力(1)、商品ID(0012)、商品名(鮭おにぎり)、単価(120)、限定割引(-)}」からV{入力(1)、商品ID(0012)、商品名(鮭おにぎり)、単価(120)、限定割引(-)}を抽出する。そして、スレーブ計算機30は、抽出したV{入力(1)、商品ID(0012)、商品名(鮭おにぎり)、単価(120)、限定割引(-)}から入力(1)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機30は、入力ファイル01用の名前付きパイプに、変換後の「商品ID(0012)、商品名(鮭おにぎり)、単価(120)、限定割引(-)」を入力する。

【0125】

同様に、スレーブ計算機30は、「K(0012)、V{入力(2)、伝票ID(0001)、商品ID(0012)、個数(2)}」からV{入力(2)、伝票ID(0001)、商品ID(0012)、個数(2)}を抽出する。そして、スレーブ計算機30は、抽出したV{入力(2)、伝票ID(0001)、商品ID(0012)、個数(2)}から入力(2)を削除し、入力ファイルと同じ形式のレコードに変換する。その後、スレーブ計算機30は、入力ファイル02用の名前付きパイプに、変換後の「伝票ID(0001)、商品ID(0012)、個数(2)」を入力する。

【0126】

10

20

30

40

50

その後、図 16 に示すように、スレーブ計算機 20 は、突合プログラムの処理結果を共有のHDFSディレクトリに出力する。具体的には、スレーブ計算機 20 は、設定ファイルで指定される「std_output」に正常処理結果を出力し、設定ファイルで指定される「error_output」に異常処理結果を出力する。なお、スレーブ計算機 30 も同様に処理結果を共有のHDFSディレクトリに出力する。

【0127】

上述したシステムは、複数の入力から突合せを行う外部プログラムに対して、MapReduceの枠組みにおいて、外部プログラムを修正なしに呼び出すことができるので、Hadoopや外部プログラムの開発生産性を向上させることができる。

【0128】

また、上述したシステムは、外部プログラムの入力ファイルをディスク上のファイルとする代わりに、名前付きパイプを用いることができるので、入力レコードの書き出しと既存プログラムの実行を並行動作させることができ、処理の高速化が図れる。また、ディスク書き込みよりも高速なプロセス間通信、すなわちメモリ間コピーを用いることができるので、スループットが向上する。また、全ての入力ファイルへの出力を同期して出力することでメモリ使用量を削減することもできる。

【0129】

また、上述したシステムは、既存プログラムの出力ファイルを回収し、MapReduceの入力として使える共有ファイルシステム上にコピーすることができる。また、上述したシステムは、外部プログラムの出力ファイルとして名前付きパイプを用い、外部プログラムがレコードを書き出すと同時に共有ファイルシステムに書き出すことができる。

【実施例 4】

【0130】

さて、これまで本発明の実施例について説明したが、本発明は上述した実施例以外にも、種々の異なる形態にて実施されてよいものである。そこで、以下に異なる実施例を説明する。

【0131】

(システム)

また、本実施例において説明した各処理のうち、自動的におこなわれるものとして説明した処理の全部または一部を手動的におこなうこともできる。あるいは、手動的におこなわれるものとして説明した処理の全部または一部を公知の方法で自動的におこなうこともできる。この他、上記文書中や図面中で示した処理手順、制御手順、具体的名称、各種のデータやパラメータを含む情報については、特記する場合を除いて任意に変更することができる。

【0132】

また、図示した各装置の各構成要素は機能概念的なものであり、必ずしも物理的に図示の如く構成されていることを要しない。すなわち、各装置の分散・統合の具体的形態は図示のものに限られない。つまり、その全部または一部を、各種の負荷や使用状況などに応じて、任意の単位で機能的または物理的に分散・統合して構成することができる。さらに、各装置にて行なわれる各処理機能は、その全部または任意の一部が、CPUおよび当該CPUにて解析実行されるプログラムにて実現され、あるいは、ワイヤードロジックによるハードウェアとして実現され得る。

【0133】

(ハードウェア構成)

ところで、上記の実施例で説明した各種の処理は、あらかじめ用意されたプログラムをパーソナルコンピュータやワークステーションなどのコンピュータシステムで実行することによって実現することができる。そこで、以下では、上記の実施例と同様の機能を有するプログラムを実行するコンピュータの一例を説明する。

【0134】

図 17 は、ハードウェア構成例を示す図である。図 17 に示すように、コンピュータ 1

10

20

30

40

50

00は、CPU101、メモリ102、ディスクドライブ103、HDD(Hard Disk Drive)104、通信制御部105、キーボード106、ディスプレイ107を有する。また、図17に示した各部は、バス100aで相互に接続される。

【0135】

通信制御部105は、NIC(Network Interface Card)などのインタフェースである。HDD104は、図2や図6等にした機能を実行するプログラムとともに、実施例1や実施例2で説明した各テーブル等を記憶する。記録媒体の例としてHDD104を例に挙げたが、ROM(Read Only Memory)、RAM、CD-ROM等の他のコンピュータが読み取り可能な記録媒体に各種プログラムを格納しておき、コンピュータに読み取らせることとしてもよい。なお、記録媒体を遠隔地に配置し、コンピュータが、その記憶媒体にアクセスすることでプログラムを取得して利用してもよい。また、その際、取得したプログラムをそのGW装置自身の記録媒体に格納して用いてもよい。

10

【0136】

CPU101は、図2に示した各処理部と同様の処理を実行するプログラムを読み出してRAMに展開することで、図2等で説明した各機能を実行するプロセスを動作させる。すなわち、このプロセスは、設定読込部14、Mapタスク管理部15、Reduceタスク管理部16を実行する。このようにコンピュータ100は、プログラムを読み出して実行することでマスタ計算機10として動作する。

【0137】

また、CPU101は、図6に示した各処理部と同様の処理を実行するプログラムを読み出してRAMに展開することで、図6等で説明した各機能を実行するプロセスを動作させる。すなわち、このプロセスは、Map初期化部24a、フォーマット初期化部24b、抽出部24c、Shuffle処理部25a、Reduce初期化部25b、再構築部25cを実行する。このようにコンピュータ100は、プログラムを読み出して実行することでスレーブ計算機20として動作する。

20

【0138】

また、コンピュータ100は、媒体読取装置によって記録媒体から上記プログラムを読み出し、読み出された上記プログラムを実行することで上記した実施例と同様の機能を実現することもできる。なお、この他の実施例でいうプログラムは、コンピュータ100によって実行されることに限定されるものではない。例えば、他のコンピュータまたはサーバがプログラムを実行する場合や、これらが協働してプログラムを実行するような場合にも、本発明を同様に適用することができる。

30

【0139】

以上の各実施例を含む実施形態に関し、さらに以下の付記を開示する。

【0140】

(付記1) コンピュータに、

形式の異なる複数の入力ファイルを読み込み、

各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成し、

前記突合キーに基づいて、各中間ファイル内のデータをソートし、

40

データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成し、

生成した複数の出力ファイルを、突合データを処理する突合プログラムに投入する

処理を実行させることを特徴とする実行制御プログラム。

【0141】

(付記2) 前記出力ファイルを生成する処理は、前記各中間ファイルのレコードごと、前記入力ファイルの各形式のデータを抽出し、

前記投入する処理は、前記各中間ファイルのレコードから前記各形式のデータが抽出されるたびに、抽出されたデータをプロセス間通信で前記突合プログラムに出力することを

50

特徴とする付記 1 に記載の実行制御プログラム。

【 0 1 4 2 】

(付記 3) 前記コンピュータは、複数のコンピュータに分散して保持されるデータを分散処理する分散処理システムを形成し、

前記突合プログラムの処理結果を、他のコンピュータが共有でアクセスできる前記分散処理システム上の共有ファイルシステムに格納する処理をさらにコンピュータに実行させることを特徴とする付記 1 または 2 に記載の実行制御プログラム。

【 0 1 4 3 】

(付記 4) 前記格納する処理は、前記突合プログラムが処理結果を出力するたびに、プロセス間通信で、前記共有ファイルシステムに処理結果を格納することを特徴とする付記 3 に記載の実行制御プログラム。

【 0 1 4 4 】

(付記 5) コンピュータが、

形式の異なる複数の入力ファイルを読み込み、

各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成し、

前記突合キーに基づいて、各中間ファイル内のデータをソートし、

データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成し、

生成した複数の出力ファイルを、データを突合する突合プログラムに入力する

を実行することを特徴とする実行制御方法。

【 0 1 4 5 】

(付記 6) 形式の異なる複数の入力ファイルを読み込む読込部と、

前記読込部によって読み込まれた各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成する第 1 生成部と、

前記突合キーに基づいて、各中間ファイル内のデータをソートするソート部と、

前記ソート部によってデータがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成する第 2 生成部と、

前記第 2 生成部によって生成された複数の出力ファイルを、データを突合する突合プログラムに入力する入力部と

を有することを特徴とする情報処理装置。

【 0 1 4 6 】

(付記 7) メモリと、

前記メモリに接続されるプロセッサと、を有し、

前記プロセッサは、

形式の異なる複数の入力ファイルを読み込み、

各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間ファイルを、前記入力ファイルごとに生成し、

前記突合キーに基づいて、各中間ファイル内のデータをソートし、

データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成し、

生成した複数の出力ファイルを、データを突合する突合プログラムに入力する

処理を実行させることを特徴とする情報処理装置。

【 0 1 4 7 】

(付記 8) 形式の異なる複数の入力ファイルを読み込み、

各入力ファイルの間で種別が共通するカラムのデータを突合キーとして追加した中間フ

10

20

30

40

50

ファイルを、前記入力ファイルごとに生成し、

前記突合キーに基づいて、各中間ファイル内のデータをソートし、

データがソートされた各中間ファイルから、前記入力ファイルの各形式のデータをそれぞれ抽出して、前記各入力ファイルに対してデータがソートされた複数の出力ファイルを生成し、

生成した複数の出力ファイルを、データを突合する突合プログラムに入力する処理をコンピュータに実行させる実行制御プログラムを記憶する、コンピュータ読み取り可能な記憶媒体。

【符号の説明】

【 0 1 4 8 】

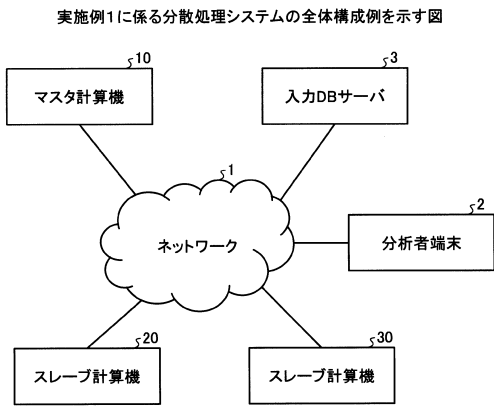
10

- 1 ネットワーク
- 2 分析者端末
- 3 入力DBサーバ
- 10 マスタ計算機
- 11 通信制御部
- 12 記憶部
- 12a 設定ファイルDB
- 12b タスクリストDB
- 13 制御部
- 14 設定読込部
- 15 Mapタスク管理部
- 16 Reduceタスク管理部
- 20、30 スレーブ計算機
- 21 通信制御部
- 22 記憶部
- 22a 中間ファイルDB
- 22b 一時ファイルDB
- 23 制御部
- 24 Map処理部
- 24a Map初期化部
- 24b フォーマット初期化部
- 24c 抽出部
- 25 Reduce処理部
- 25a Shuffle処理部
- 25b Reduce初期化部
- 25c 再構築部

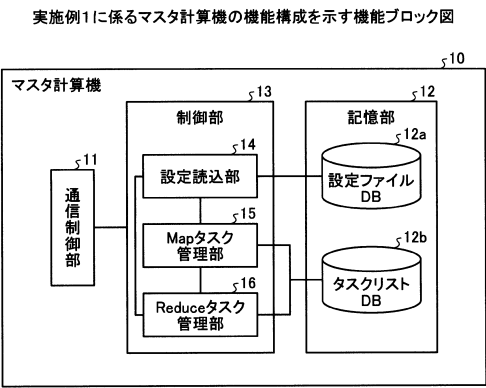
20

30

【図 1】



【図 2】



【図 3】

設定ファイルDBに記憶される情報の例を示す図

ジョブ設定
入力データファイル名01= /usr/transaction.txt
入力データファイル名02= /usr/master.txt
入力フォーマット処理クラス名01= 固定長フォーマット
入力フォーマット処理オプション01= レコード長:80BYTE
入力フォーマット処理クラス名02= 改行ありフォーマット
入力フォーマット処理オプション02= 改行コード:¥r¥n
Reduceアプリケーション= a.out %in01 %in02 %out01 %out02
入力ファイル環境変数名01= TRAN
入力ファイル環境変数名02= MASTAR
出力データディレクトリ名= /output/
出力ファイル名01= processed.txt
出力ファイル名02= error.txt
出力ファイル環境変数名01= PROCOUT
出力ファイル環境変数名02= ERROUT
ソートキー01= 4BYTE~8BYTE
ソートキー02= 12BYTE~16BYTE

【図 4】

タスクリストDBに記憶される情報の例を示す図

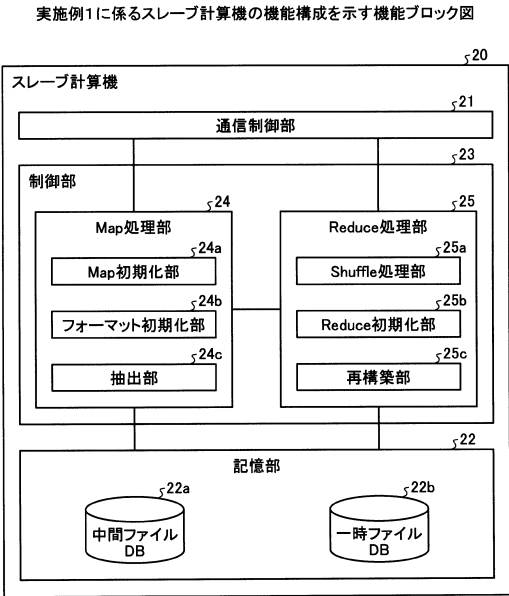
タスクID	入力情報	クラス名
task_m_1	ファイル名=/usr/transaction.txt	フォーマット処理=本フォーマット処理01
	開始位置=0 サイズ=1000	マップ処理=本Map処理01
task_m_2	ファイル名=/usr/transaction.txt	フォーマット処理=本フォーマット処理01
	開始位置=1001 サイズ=1000	マップ処理=本Map処理01
task_m_3	ファイル名=/usr/transaction.txt	フォーマット処理=本フォーマット処理01
	開始位置=2001 サイズ=1000	マップ処理=本Map処理01
task_m_4	ファイル名=/usr/master.txt	フォーマット処理=本フォーマット処理02
	開始位置=0 サイズ=1000	マップ処理=本Map処理02
task_m_5	ファイル名=/usr/master.txt	フォーマット処理=本フォーマット処理02
	開始位置=1001 サイズ=1000	マップ処理=本Map処理02

【図 5】

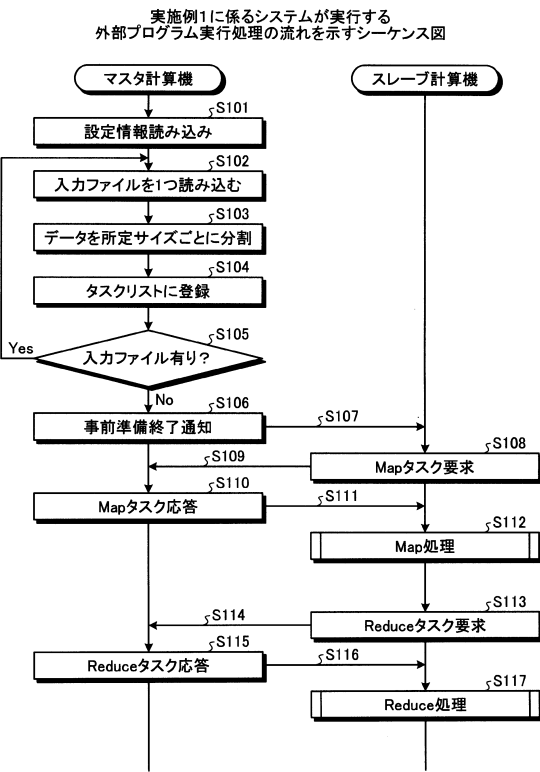
スレーブ計算機に送信されるReduceタスク情報の例を示す図

タスクID	Map結果情報		クラス名
	MapタスクID	スレーブ計算機名	
task_r_1	task_m_1	スレーブ計算機20	本Reduceクラス
	task_m_2	スレーブ計算機30	

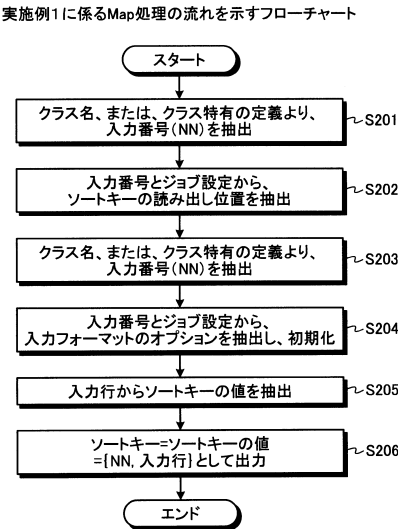
【図 6】



【図 7】



【図 8】



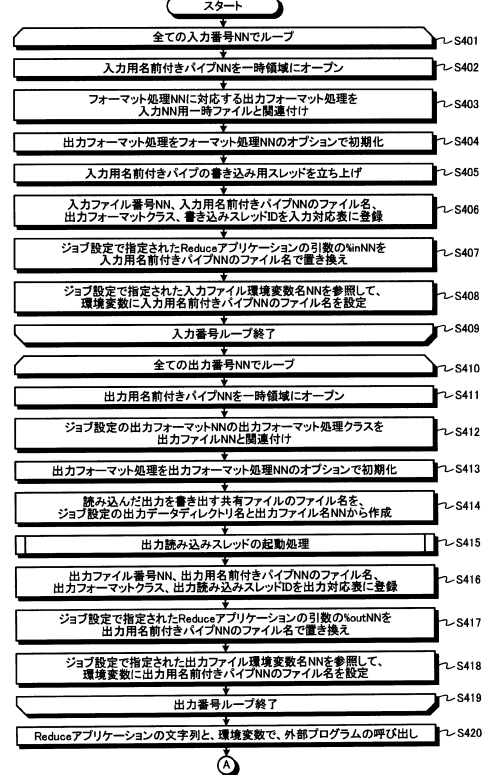
【図 9】

実施例1に係るReduce処理の流れを示すフローチャート

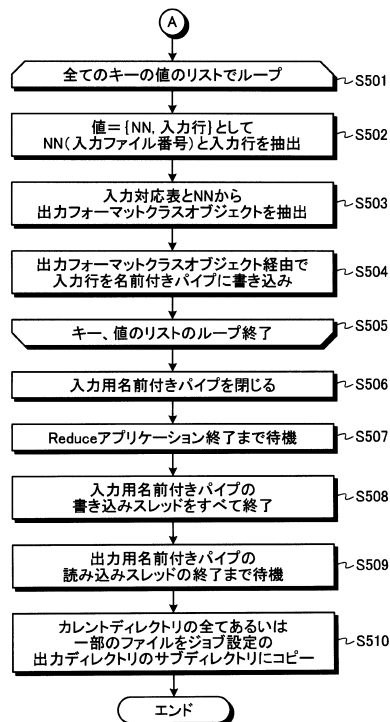


【図 10】

実施例2に係るReduce処理の初期化の流れを示すフローチャート

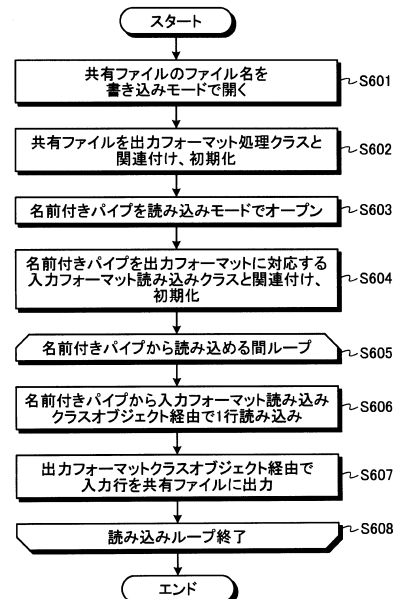


【図 11】

実施例2に係るReduce処理の
本処理および完了処理の流れを示すフローチャート

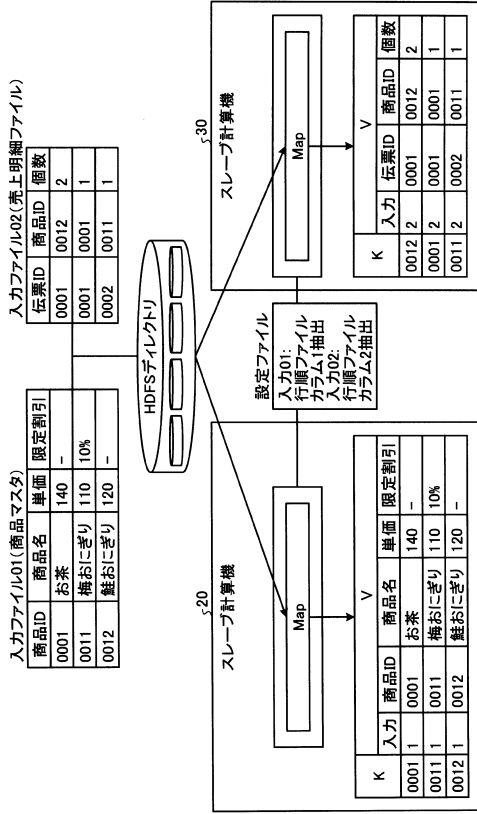
【図 12】

実施例2に係る出力読み込みスレッドの起動処理の流れを示すフローチャート



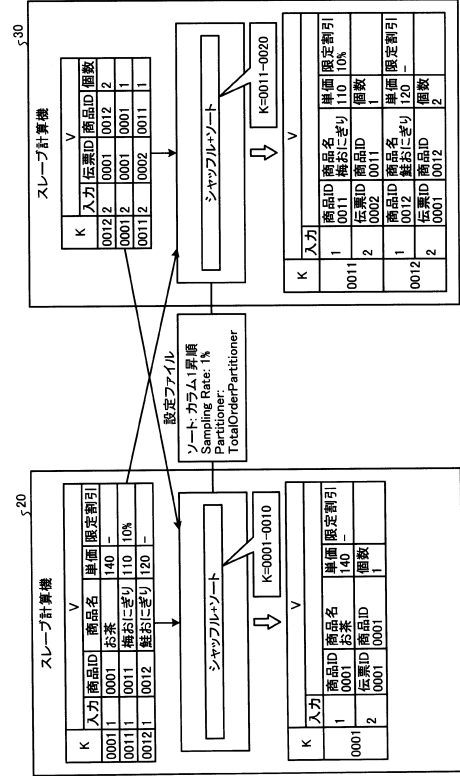
【図 13】

Map処理の具体例を説明する図



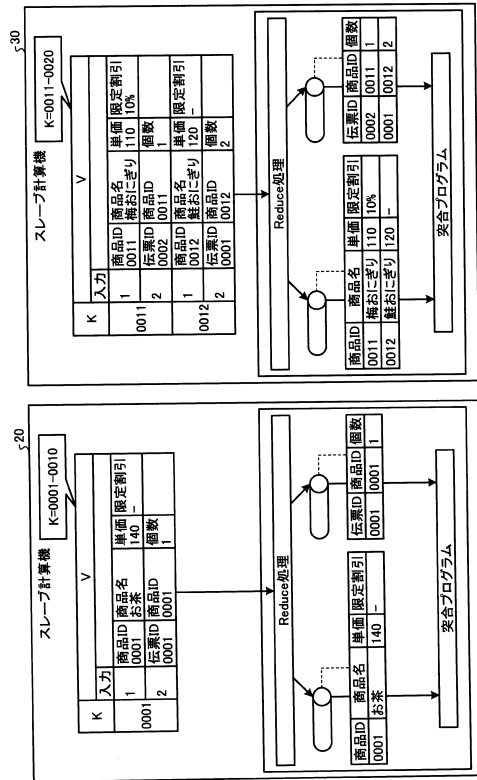
【図 14】

シャッフルソート処理の具体例を説明する図



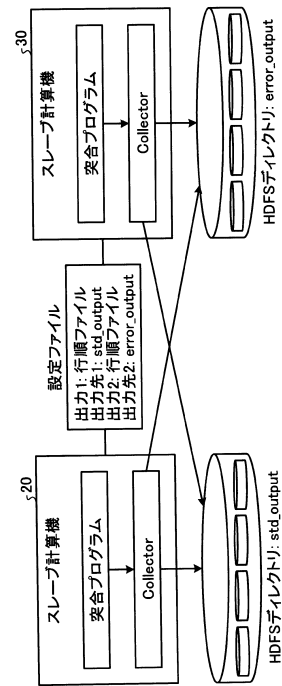
【図 15】

Reduce処理の具体例を説明する図



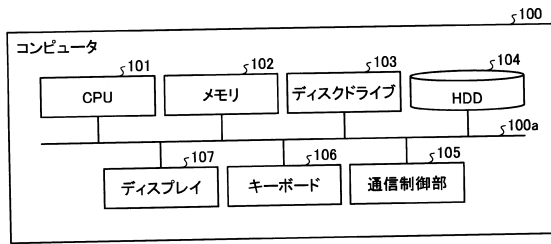
【図 16】

実行プログラム実行後の出力先の具体例を説明する図



【図 17】

ハードウェア構成例を示す図



フロントページの続き

(58)調査した分野(Int.Cl. , D B 名)

G 0 6 F 1 7 / 3 0