

(19) 日本国特許庁(JP)

(12) 公表特許公報(A)

(11) 特許出願公表番号

特表2016-534471

(P2016-534471A)

(43) 公表日 平成28年11月4日(2016.11.4)

(51) Int.Cl.	F I	テーマコード (参考)
G06F 11/14 (2006.01)	G06F 11/14 669	5B027
G06F 13/10 (2006.01)	G06F 13/10 340A	
G06F 3/06 (2006.01)	G06F 3/06 304F	
G06F 12/00 (2006.01)	G06F 3/06 306Z	
	G06F 11/14 664	
審査請求 有 予備審査請求 未請求 (全 53 頁) 最終頁に続く		

(21) 出願番号 特願2016-540866 (P2016-540866)
 (86) (22) 出願日 平成25年10月18日 (2013.10.18)
 (85) 翻訳文提出日 平成28年3月2日 (2016.3.2)
 (86) 国際出願番号 PCT/US2013/065623
 (87) 国際公開番号 W02015/057240
 (87) 国際公開日 平成27年4月23日 (2015.4.23)

(71) 出願人 505160898
 ヒタチ データ システムズ エンジニア
 リング ユーケー リミテッド
 HITACHI DATA SYSTEM
 S ENGINEERING UK LI
 MITED
 イギリス国 アールジー12 8エフゼッ
 ト ブラックネル, オールドベリー, キャ
 ピトル ビルディング
 (74) 代理人 110000279
 特許業務法人ウィルフォート国際特許事務
 所

最終頁に続く

(54) 【発明の名称】 シェアード・ナッシング分散型ストレージ・システムにおけるターゲットにより駆動される独立したデータの完全性および冗長性のリカバリ

(57) 【要約】

ネットワークに接続された独立したノードと、ストレージ・システムに記憶されたデータに対するI/O処理を送るホストとを含む分散型シェアード・ナッシング・ストレージ・システム。各ノードは、部分的に書き込まれたデータのそのノード自体のリカバリを実行し、ストレージ・システムに記憶されたデータの整合性を維持することができる。ノードが、ノード中のすべてのデータの位置を独立して計算し、独立してデータの釣り合いを取り、ストレージ・システムの冗長性のポリシーに基づいて整合性を維持し、位置の変更に応じてデータをマイグレーションする。ノードは、そのノードが不完全なまたは損傷したデータを記憶すると判定する場合、その他のノード上のレプリカ・データからそのノードが記憶するデータを再構築する。ノードの間のデータのマイグレーション中、ホストからのI/O処理は、妨げられない。

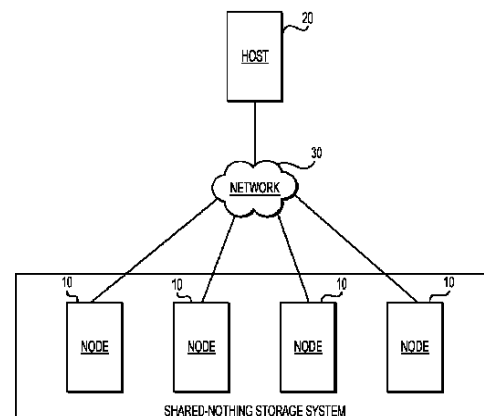


FIG. 1

【特許請求の範囲】**【請求項 1】**

1 つまたは複数のネットワークに接続されたクラスタ内の複数の独立したノードであって、それぞれが、データを記憶するための 1 つまたは複数のデバイスを含み、

それぞれの独立したノードが、入力 / 出力 (I / O) オペレーションを送り、受け取るように構成され、

前記クラスタ内の前記複数のノードに記憶された前記データが、前記データレプリカを含み、各ノードが、位置のマップに基づいて前記クラスタ内のストレージ・デバイス上に記憶された前記データおよび前記データレプリカの位置を知っている、複数の独立したノードと、

10

前記クラスタ内のすべてのノードが利用可能である、各ノードおよび各ノードの各ストレージ・デバイスの状態を示す状態情報とを含み、

前記クラスタ内でノードがアクティブ化もしくは非アクティブ化されるかまたはいずれかのノードのストレージ・デバイスがアクティブ化もしくは非アクティブ化されると、前記状態情報が、対応するがアクティブ化もしくは非アクティブ化に基づいて更新され、変更され、

前記状態情報の変更に基づいて、新しくアクティブ化されたノードまたは新しくアクティブ化されたストレージ・デバイスを有するノードを含む前記クラスタ内の各ノードが、前記クラスタ内の前記ノード上に記憶されるいずれかのデータまたはデータレプリカの整合性を維持するためにリカバリオペレーションを前記ノードが実行しなければならないかどうかを判定し、

20

各ノードが、影響を与えられるデータを記憶する前記クラスタ内のその他のノードにメッセージを送信し、受信することによって、前記クラスタ内のその他のノードとは独立して必要に応じて前記リカバリオペレーションを実行して、前記ノードのストレージ・デバイスに記憶された前記データを復元する分散型シェアード・ナッシング・ストレージ・システム。

【請求項 2】

前記クラスタ内のデータが、前記クラスタ内の複数のノードにレプリカまたは消失訂正符号化されたセグメントとして冗長的に記憶され、

レプリカまたは消失訂正符号化されたセグメントの数が、前記データの冗長性のポリシーに依存し、

30

レプリカまたは消失訂正符号化されたセグメントの符号化が、前記データの冗長性のポリシーおよび消失訂正符号化アルゴリズムに依存する請求項 1 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 3】

各ノードが、イベントの通知を受信し、前記イベントが、前記クラスタ内のノードの前記アクティブ化もしくは前記非アクティブ化またはいずれかのノードのストレージ・デバイスの前記アクティブ化もしくは前記非アクティブ化である請求項 2 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 4】

40

各イベントが、厳密な順番で前記クラスタ内の各ノードに届けられ、前記イベントが、単調な昇順で一意で整合性のある番号を振られ、

各ノードが、前記クラスタ内のすべてのイベントを受信することを保証され、

非アクティブ化の後でアクティブ化された各ノードが、前記ノードがアクティブ化されるときにすべてのイベントを受信することを保証される請求項 1 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 5】

前記状態情報が、各ノードのストレージ・デバイスに記憶されたデータおよび各ノードのストレージ・デバイス上のデータレプリカまたは消失訂正符号化されたセグメントの状態、ならびに前記クラスタ内のすべての前記その他のノードのストレージ・デバイス上の

50

データおよびデータレプリカまたは消失訂正符号化されたセグメントの状態についての情報を含み、

前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントの前記状態が、

前記データを修正または書き込みする前にデータ、データレプリカ、または消失訂正符号化されたセグメントに関して設定される D I R T Y フラグ、

データの書き込みまたは修正がデータレプリカまたは消失訂正符号化されたセグメントの単に一部 (less than all) で成功したことを前記ノードが知っているときにデータ、前記データレプリカ、または前記消失訂正符号化されたセグメントに関して設定される D E G R A D E D フラグ、

イベント番号の後に前記データが書き込まれた前記イベント番号を示す情報、およびイベント番号の後に前記 D E G R A D E D フラグが設定された前記イベント番号を示す情報を含み、

前記 D I R T Y フラグが永続的であり、前記 D E G R A D E D フラグが永続的であり、

前記 D I R T Y フラグが、前記データのデータレプリカまたは消失訂正符号化されたセグメントの少なくとも予め決められた最小限の数に対する書き込みが正常に完了したときにクリアされ、

前記データおよび前記データレプリカまたは前記消失訂正符号化されたセグメントの位置が、永続的に記憶される可能性があり、または前記データの前記位置のマップおよびメタデータから必要に応じて計算される可能性があり、

前記メタデータが、一意の I D、前記データの名前空間内の一意の名前属性、または前記データの内容の記述のいずれかを含み請求項 3 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 6】

前記クラスタ内の各ノードによって受信される前記イベントが、非アクティブ化の後にアクティブ化された第 1 のノードのアクティブ化または第 1 のノードのストレージ・デバイスのアクティブ化であり、

前記判定が、前記位置のマップに基づいていずれかのノードが前記第 1 のノードに記憶されたいずれかのデータ、データレプリカ、または消失訂正符号化されたセグメントを有するかどうかに基づき、

前記第 1 のノードが、前記 D E G R A D E D フラグ、前記 D I R T Y フラグ、前記データの最後の修正または書き込みのイベント番号を示す情報、およびイベント番号の後に前記 D E G R A D E D フラグが設定された前記イベント番号を示す前記情報を含む情報のリストを受信し、

情報の前記リストに基づいて、前記第 1 のノードが、どのデータ、レプリカ、または消失訂正符号化されたセグメントが再構築される必要があるかを判定し、

前記状態情報に基づいて、前記第 1 のノードが、前記第 1 のノードのストレージ・デバイス上の前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントを再構築するために再構築される必要がある前記第 1 のノード上のデータの前記レプリカまたは前記消失訂正符号化されたセグメントを前記クラスタ内のどのノードが記憶するかを判定し、

前記第 1 のノードが、再構築される必要がある前記データを記憶すると判定された前記クラスタ内の前記ノードからの、再構築される前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントの内容を要求し、再構築すべき前記データの前記内容を書き込み、完了すると、情報の前記リストを送信したすべてのノードにメッセージを送信し、

前記データのすべてのデータレプリカまたは消失訂正符号化されたセグメントが前記クラスタ内の各ノードでリカバリされると判定される場合に、前記クラスタ内のすべてのノードが、前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントの前記 D E G R A D E D フラグをクリアする請求項 5 に記載の分散型シェアード・ナッシ

10

20

30

40

50

ング・ストレージ・システム。

【請求項 7】

ネットワークに接続されたクラスタ内の複数の独立したノードであって、それぞれが、データを記憶するための 1 つまたは複数のストレージ・デバイスを含み、

それぞれのノードが、入力 / 出力 (I / O) オペレーションを送り、受け取るように構成され、

各ノードが、ローカルの処理を行うための、前記ストレージ・デバイスとは別であるメモリを有する、複数の独立したノードを含み、

I / O オペレーションが、個々のノードで実行され、I / O オペレーションを受け取る各ノードが、前記ノードのストレージ・デバイスに記憶された前記データに関してのみ前記 I / O オペレーションを実行し、

前記クラスタ内の前記ストレージ・デバイスに記憶されたデータの位置が、すべての前記ノードが利用可能であるレイアウト情報から導出され、

前記データの前記位置が、記憶されるか、前記レイアウト情報から必要に応じて各ノードによって計算され、

イベントの発生を示す第 1 のイベント・メッセージを受信すると、各ノードが、前記ノードのストレージ・デバイス上に記憶されたすべてのデータを評価し、いずれかのデータが前記第 1 のイベント・メッセージの情報によって影響を与えられるかどうかを判定する分散型シェアード・ナッシング・ストレージ・システム。

【請求項 8】

各ノードおよび各ノードの各ストレージ・デバイスがアクティブ状態であるかどうかを示す、すべてのノードが利用可能な状態情報をさらに含み、

レイアウト情報が、

前記ストレージ・システム内の各ノードおよびストレージ・デバイス上のデータの位置を決定し、I / O 処理のためにデータの前記位置を判定するために各ノードによって使用される現在のレイアウト・マップと、

前記ストレージ・システム内の各ノードおよびストレージ・デバイス上のデータの提案される位置を決定し、I / O オペレーションのためにデータの前記位置を決定するために各ノードによって使用され、前記第 1 のイベント・メッセージを受信すると各ノードにおいて生成される提案されるレイアウト・マップであって、前記現在のレイアウト・マップと比較して、前記イベントに基づいて少なくとも 1 つのデータの位置を変更する、提案されるレイアウト・マップとを含み、

前記第 1 のイベント・メッセージが、前記イベントが発生すると現在のマップ・レイアウトの下で前記クラスタ内のすべてのノードに決められた順番に送信され、前記第 1 のイベント・メッセージの前記情報が、イベントの種類を示し、

各ノードが、メモリに、前記現在のレイアウト・マップおよび前記提案されるレイアウト・マップを記憶し、

すべてのノードが第 2 のイベント・メッセージを受信すると、前記ノードが、現在のレイアウト・マップの下から前記提案されるレイアウト・マップの下へのデータのマイグレーションを開始し、

前記現在のレイアウト・マップが、第 3 のイベント・メッセージを受信されると、すべての前記ノードによって前記提案されるレイアウト・マップに置き換えられ、

前記現在のレイアウト・マップおよび前記提案されるレイアウト・マップの下の対応するノードが、マイグレーションされている前記データに関する I / O オペレーションを受け付け続ける請求項 7 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 9】

前記イベントが、前記クラスタ内のノードのアクティブ化もしくは非アクティブ化またはいずれかのノードのストレージ・デバイスのアクティブ化もしくは非アクティブ化であり、

前記第 1 のイベント・メッセージを受信されると、各ノードが、前記ノードがソース・

10

20

30

40

50

ノードであるかどうかを判定し、各ソース・ノードが、前記提案されるレイアウト・マップの下でのデータの前記位置のノードに、マイグレーション要求を送信することによって、前記ノードが前記データの内容を受信すべきであることを知らせ、

マイグレーション要求を受信する各ノードが、前記提案されるレイアウト・マップの下で前記ノードが記憶する前記データの前記内容を獲得するリカバリ要求を送信する請求項 8 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 10】

前記クラスタ内のデータが、前記クラスタ内の複数のノードにレプリカまたは消失訂正符号化されたセグメントとして冗長的に記憶され、

レプリカまたは消失訂正符号化されたセグメントの数が、前記データの冗長性のポリシーに依存し、

前記レプリカまたは前記消失訂正符号化されたセグメントの位置が、前記現在のレイアウト・マップおよび前記提案されるレイアウト・マップに含まれ、

レプリカまたは消失訂正符号化されたセグメントの符号化が、前記ストレージ・システムの前記冗長性のポリシーおよび消失訂正符号化アルゴリズムに依存し、

前記データレプリカおよび消失訂正符号化されたセグメントの位置が、前記レイアウト情報から導出される請求項 9 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 11】

データに関する I/O オペレーションが、前記現在のマップ・レイアウトと提案されるマップ・レイアウトとの両方の下で前記データに関する前記レプリカまたは前記消失訂正符号化されたセグメントを記憶するすべてのノードに送られ、

前記提案されるマップ・レイアウトの下でのノードが、I/O オペレーションが影響を与えるデータを記憶する前記現在のマップ・レイアウトの下でのノードから前記データに関する前記 I/O オペレーションを受け取る請求項 10 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 12】

すべてのノードが第 2 のイベント・メッセージを受信した後、I/O オペレーションを受け取る前記提案されるマップ・レイアウトの下でのノードが、前記 I/O オペレーションが影響を与える前記データが前記提案されるマップ・レイアウトの下での前記ノードにマイグレーションされたか否かを判定し、

前記 I/O オペレーションが影響を与える前記データがマイグレーションされていない場合、前記提案されるマップ・レイアウトの下での前記ノードが、前記オペレーションを実行するために前記データがマイグレーションされるまで待ち、

前記 I/O オペレーションが影響を与える前記データがマイグレーションされた場合、前記提案されるマップ・レイアウトの下での前記ノードが、前記データに対して前記 I/O オペレーションを実行する請求項 11 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 13】

前記第 1 のイベント・メッセージによって影響を与えられる判定された各データに関して、前記ノードが、前記ノードが前記現在のマップ・レイアウトの下での前記データのコピーを有するが、前記提案されるマップ・レイアウトの下での前記データのコピーを持たないかどうか、または前記ノードが前記現在のマップ・レイアウトと前記提案されるマップ・レイアウトとの両方下のデータを記憶するかどうか、または前記ノードが前記提案されるレイアウト・マップの下でのノードに、前記ノードが記憶すべきデータを有し、前記データが前記位置にマイグレーションされるべきであることを知らせる役割を担うノードであるかどうかを判断し、

前記ノードが前記提案されるマップ・レイアウトの下でのノードに知らせる役割を担うかどうかの判定が、レイアウト情報および状態情報に基づく請求項 12 に記載の分散型シェアード・ナッシング・ストレージ・システム。

10

20

30

40

50

【請求項 14】

前記データが、データのサブセットであり、有限のサイズであるスライスでマイグレーションされ、

前記 I/O オペレーションが W R I T E 要求である場合、および前記 W R I T E 要求が影響を与えるデータがマイグレーションされていない場合、前記提案されるマップ・レイアウトの下の前記ノードが、前記 W R I T E 要求のサイズがスライスのサイズ未満であるかどうかを判定し、

前記 W R I T E 要求の前記サイズがスライスの前記サイズ未満である場合、前記 W R I T E 要求が当てはまる前記データのサイズが直ちにマイグレーションされ、

前記 W R I T E 要求の前記サイズがスライスの前記サイズ未満でないか、または前記スライスの前記サイズに等しい場合、前記 W R I T E 要求に関して書き込まれるデータが、マイグレーションされるスライスに関して書き込まれ、前記スライスが、マイグレーションされない請求項 12 に記載の分散型シェアード・ナッシング・ストレージ・システム。

10

【請求項 15】

分散型シェアード・ナッシング・ストレージ・システムであって、

ネットワークに接続されたクラスタ内の複数の独立したノードであって、それぞれが、データを記憶するための 1 つまたは複数のストレージ・デバイスを含み、

それぞれのノードが、入力/出力 (I/O) オペレーションを送り、受け取るように構成される、複数の独立したノードと、

各ノードおよび各ノードの各ストレージ・デバイスがアクティブ状態であるかどうかを示す、すべてのノードが利用可能な状態情報とを含み、

20

前記クラスタ内のデータが、前記クラスタ内の複数のノードにレプリカまたは消失訂正符号化されたセグメントとして冗長的に記憶され、

レプリカまたは消失訂正符号化されたセグメントの数が、前記データの冗長性のポリシーに依存し、

レプリカまたは消失訂正符号化されたセグメントの符号化が、前記データの冗長性のポリシーおよび消失訂正符号化アルゴリズムに依存し、

前記クラスタ内の前記ノードの前記ストレージ・デバイス上の前記データおよび前記データのレプリカまたは前記消失訂正符号化されたセグメントの位置が、すべての前記ノードが利用可能であるレイアウト情報から導出され、

30

前記データの位置が、記憶されるか、前記レイアウト情報から必要に応じて計算され、

ストレージ・システムに関する冗長性のポリシーが、前記データのデータレプリカもしくは消失訂正符号化されたセグメントの数を示すミラー・カウント数を設定し、データレプリカもしくは消失訂正符号化されたセグメントが、前記クラスタ内の前記ストレージ・デバイス上の第 1 の位置から第 2 の位置にマイグレーションされ、またはデータ、データレプリカ、もしくは消失訂正符号化されたセグメントが、ストレージ・システムの前記冗長性のポリシーの変更に基づいて前記クラスタから削除される、分散型シェアード・ナッシング・ストレージ・システム。

【請求項 16】

前記レイアウト情報が、

40

前記ストレージ・システム内の各ノードおよびストレージ・デバイス上のデータの位置を決定し、データ、データレプリカ、または消失訂正符号化されたセグメントの前記位置を判定するために各ノードによって使用される現在のレイアウト・マップと、

前記ストレージ・システム内の各ノードおよびストレージ・デバイス上のデータの提案される位置を決定し、データ、データレプリカ、または消失訂正符号化されたセグメントの前記位置を判定するために各ノードによって使用され、前記現在のレイアウト・マップの下各ノードが第 1 のポリシー変更イベント・メッセージを受信すると各ノードで生成される提案されるレイアウト・マップとを含み、前記提案されるレイアウト・マップが、前記現在のレイアウト・マップと比較してポリシー・イベント・メッセージに含まれる前記冗長性のポリシーの変更に基づいて少なくとも 1 つのデータ、データレプリカ、または

50

消失訂正符号化されたセグメントの前記位置を変更する請求項 15 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 17】

前記冗長性のポリシーの前記変更が、前記ミラー・カウントの増加であるかまたは前記ミラー・カウントの減少かのどちらかであり、

前記変更がミラー・カウントの減少である場合、前記変更が影響を与える前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントを記憶する各ノードが、前記ノードが前記提案されるレイアウト・マップの下の前記データを記憶するかどうかを判定し、前記ノードが前記提案されるレイアウト・マップの下の前記データ、前記レプリカ、または前記消失訂正符号化されたセグメントを記憶しない場合、前記ノードが、前記ノードのストレージ・デバイスから前記データ、前記レプリカ、または前記消失訂正符号化されたセグメントを削除し、

前記変更がミラー・カウントの増加である場合、前記提案されるレイアウト・マップが、前記データ、前記レプリカ、または前記消失訂正符号化されたセグメントの追加のミラーに関する新しい位置を決定し、

前記変更が影響を与える各データ、レプリカ、消失訂正符号化されたセグメントに関して、各ノードが、前記新しい位置のノードに、マイグレーション要求を送信することによって、前記ノードが前記変更が影響を与える前記データ、前記レプリカ、または前記消失訂正符号化されたセグメントの内容を受信すべきであることを知らせる請求項 16 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 18】

マイグレーション要求を受信する各ノードが、前記提案されるレイアウト・マップの下で前記ノードが記憶する前記データ、前記データレプリカ、または消失訂正符号化されたセグメントの前記内容を獲得するリカバリ要求を送信し、

前記マイグレーションが完了すると、前記提案されるレイアウト・マップが前記現在のレイアウト・マップを置き換える請求項 17 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 19】

前記データ、前記データレプリカ、または前記消失訂正符号化されたセグメントの前記新しい位置が、前記現在のレイアウト・マップに基づく前記データ、前記レプリカ、または前記消失訂正符号化されたセグメントの位置と前記提案されるレイアウト・マップに基づく前記データの位置とを比較することによって各ノードにより計算される請求項 17 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【請求項 20】

前記現在のレイアウト・マップおよび前記提案されるレイアウト・マップの下の対応するノードが、マイグレーションされている前記データに関する I/O オペレーションを受け付け続け、

データに関する I/O オペレーションが、現在のマップ・レイアウトと提案されるマップ・レイアウトとの両方の下で前記データに関する前記レプリカまたは前記消失訂正符号化されたセグメントを記憶するすべてのノードに送られ、

前記提案されるマップ・レイアウトの下で前記ノードが、I/O オペレーションが影響を与えるデータを記憶する前記現在のマップ・レイアウトの下で前記ノードから前記データに関する前記 I/O オペレーションを受け取る請求項 15 に記載の分散型シェアード・ナッシング・ストレージ・システム。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、ネットワークによって接続された独立したコンピュータ・システム（ノード）のクラスタからなるシェアード・ナッシング分散型ストレージ・システムに関する。

【背景技術】

【 0 0 0 2 】

ストレージ・システムに記憶されるデータは、高い信頼性で保有されなければならない、入力／出力（I／O）処理は、ストレージ・システムにおけるデータのマイグレーション中に妨げられてはならない。たとえば、書き込みオペレーションに関して、マイグレーション中に、ストレージ・システムは、データ・オブジェクトの状態を高い信頼性で追跡しなければならない。

【 0 0 0 3 】

1つの知られている方法は、修正されるべきデータの領域が、たとえば、データを書き込む前に共通の／共有された「スコアボード（scoreboard）」上で「ダーティ」フラグによって印を付けられる書き込みのマーキングを使用する。この手法では、データを
10
書き込む要求をログに記録するステップと、データを記憶する各ターゲットにメッセージを送信するステップと、書き込みおよび応答を待つステップと、それから、実際の書き込みオペレーションを送るステップとを含むいくつかのステップが、必要とされる。上述の方法は、書き込みオペレーションに関するネットワーク・レイテンシの増加につながる。

【 0 0 0 4 】

別の知られている記憶方法は、高レベルのデータ・ストレージ・エリア全体にダーティとして印を付ける。しかし、そのような手法は、データの大きな集合体全体のリカバリを必要とするので、大量のデータに対して実用的ではない。知られているストレージ・システムは、修正を示すためにファイル・システム・レベルでファイルにダーティとして印を
20
付ける可能性もある。しかし、ファイル・システム・レベルのマーキングは、非常に大きなデータ・ファイルに関して有効であり得ないほど粗い粒度を有する印を付けられたデータをもたらし、その結果、完了し得ないほど長い期間を必要とするリカバリをもたらす。さらに、Parascale Inc. のスケールアウト・ストレージ・プラットフォーム・ソフトウェアなど、集中型のデータベースにおいてデータのチャンクにダーティとして印を付けることも、当技術分野で知られている。

【 0 0 0 5 】

知られているストレージ・システムにおける同様の機能は、たとえば、参照により本明細書に援用される米国特許第6,907,507号、第6,910,111号、第7,089,385号、および第6,978,354号に記載されているVERITAS Volume Manager (VxVM) のFast Mirror Resync (FMR) の特徴をさらに含む。米国特許第6,907,507号、第6,910,111号、第7,089,385号、および第6,978,354号は、多カラムのビットマップ、
30
アキュムレータのマップ、およびミラー毎のマップを使用する。I／Oエラーからのリカバリに関して、従来技術のストレージ・システム（ボリューム・マネージャおよび多コピー・ファイル・システム）は、データを直接読み取るかもしくは書き込むかのどちらかによってリカバリを実行するための中央マネージャを必要とするか、またはリカバリ・プロセスを管理するためのコーディネータを必要とする。そのような構成の欠点は、コーディネータが障害にあうときに集中的に管理されたリカバリが止まり、それがリカバリ・プロセスのさらに困難な状況につながることである。したがって、コーディネータの障害の
40
可能性を考慮して、大量のメタデータが共有されたストレージに高い信頼性で保有されることが求められる。

【 0 0 0 6 】

部分的に書き込まれたデータのリカバリの場合、従来技術は、多くのボリューム・マネージャの実装で使用するミラー再接続およびミラー「リシルバリング（resilve
ring）」手法からなり、それらの手法は、中央データベースまたは何らかの種類の
50
ボリューム・レベルのビットマップを使用する。その他の実装は、1つの中央の位置（すべてのボリューム・マネージャ）からの直接読み取りおよび書き込みを行う中央リカバリ・マネージャを使用するか、または（たとえば、Parascale Inc. のスケールアウト・ストレージ・プラットフォーム・ソフトウェアのように）リカバリを駆動するた

めの中央コーディネータを有する。

【0007】

ノードまたはノードのディスクがストレージ・システムにおいて追加または削除されるデータのマイグレーションを伴う場合、従来技術は、信頼性の高いハッシュおよびマップの生成に基づくCEPHファイル・システムの再レイアウトの特徴を含む。PICASSOシステムとCEPHシステムとの両方は、「CRUSH」アルゴリズムとしてよく知られている配置アルゴリズムを用いて、ストレージ・クラスタ全体にまたがるストレージ構成のバージョン情報に基づいてデータ・チャンクの適切な配置を決定的に計算する。参照により本明細書に援用されるSage A. Welli、Scott A. Brandt、Ethan L. Miller、Carlos Maltzahn、「CRUSH: Controlled, Scalable, Decentralized Placement of Replicated Data」、Proceedings of the 2006 ACM/IEEE Conference on Supercomputing、p. 31、2006を参照されたい。CEPHシステムにおいて、再レイアウトは、中央メタデータ・エンジンによって実行される。さらに、Parascaleシステムにおいて、データの再レイアウトは、中央データベースによって駆動され、配置は、アドホックにチャンク毎に行われる。Parascaleシステムにおける再レイアウトが妨げられるとき、データのレイアウトは過渡的であるが整合性のある状態で残され、再レイアウト・プロセスが再開すると、データの配置が再計算される。

10

20

【先行技術文献】

【特許文献】

【0008】

【特許文献1】米国特許第6,907,507号

【特許文献2】米国特許第6,910,111号

【特許文献3】米国特許第7,089,385号

【特許文献4】米国特許第6,978,354号

【非特許文献】

【0009】

【非特許文献1】Sage A. Welli、Scott A. Brandt、Ethan L. Miller、Carlos Maltzahn、「CRUSH: Controlled, Scalable, Decentralized Placement of Replicated Data」、Proceedings of the 2006 ACM/IEEE Conference on Supercomputing、p. 31、2006

30

【非特許文献2】Leslie Lamportによる「Paxos Made Simple」、ACM SIGACT News (Distributed Computing Column) 32、4 (通巻号数121、2001年12月) 51~58

【非特許文献3】Tushar Deepak Chandra、Robert Griesemer、およびJoshua Redstoneによる「Paxos Made Live - An Engineering Perspective (2006 Invited Talk)」、Proceedings of the 26th Annual ACM Symposium on Principles of Distributed Computing、ACM press (2007)

40

【発明の概要】

【発明が解決しようとする課題】

【0010】

データの冗長性に対してポリシーの変更が行われる場合、データの移動は、集中的に管理され、中央管理ノードから実行される。Parascale Inc. のシステムは、新しいデータの位置として決定される位置が冗長性の変更を満たすために既存の記憶位置からデータを「引き出す」ことを求められるマイグレーションの中央管理を行っていた。

50

【課題を解決するための手段】

【0011】

本発明の実施形態は、概して、前記コンピュータ・システムを制御するコンピュータ・システムおよび対応する方法に関し、より詳細には、シェアード・ナッシング分散型ストレージ・システムにおけるデータ管理およびマイグレーションのための技術に関する。下で説明される以下の実施形態は、例示的であり、本発明が本発明の精神および範囲内に入るその他の修正、代替的な構造、および均等物によって実施される可能性があることと、さらに、本発明の1つまたは複数の実施形態と一緒に組み合わせられる可能性があることとを当業者が理解するような性質を持つ。

【0012】

一実施形態においては、シェアード・ナッシング分散型ストレージ・システムにおいてデータの完全性および冗長性を保ち、データの完全性および冗長性をリカバリする方法が提供される。シェアード・ナッシング分散型ストレージ・システムにおいて、ストレージ・システムは、ストレージ・システム（本明細書においてはクラスタと呼ばれる）を形成するために相互に接続される複数の独立したシステム（本明細書においてはノードと呼ばれる）を含む。クラスタの各ノードは、ファイルの名前空間、メタデータ、および位置情報の管理を扱う。さらに、各ノードは、クラスタ内の他のノードと独立しているストレージ・エリアを備え、クラスタの1つまたは複数のファイル・システムに記憶されるファイルへの物理的な記憶およびアクセスの責任を負う。ローカルであるかまたはストレージ・エリア・ネットワーク（SAN）などを通じて取り付けられるかに関わらず、各ストレージ・ノードは1つまたは複数の永続的なストレージ・デバイスを備えており、通常の場合で、各ノードは、そのノード自体のストレージ・デバイスにのみアクセスすることができ、概して、ノードは、その他のノードのストレージ・デバイスにアクセスすることができず、したがって、システム全体のストレージが、システムのノードの間に分散されるので、このクラスタ・アーキテクチャは、シェアード・ナッシング分散型ストレージ・システムと呼ばれる。

【0013】

データは、ストレージのステータスおよびクラスタ全体の活動の情報を考慮に入れるアルゴリズムを用いて、または何らかのその他のマッピング機能を用いて、クラスタ内のストレージ・デバイスに分散される可能性がある。例示的な実施形態においては、トポロジーおよびクラスタ内のストレージの負荷要因を用いて、ならびにデータに関連する一意の「キー」を用いてデータが配置される位置を計算してクラスタ内のデータの配置を決定する擬似ランダムなコンシステント・ハッシュ（CRUSHアルゴリズム）が使用される。例示的な実施形態において、データは、別個の有限なサイズのコンテナ（本明細書においては「チャンク」と呼ばれる）にマッピングされる。これらのチャンクは、チャンクまたはチャンクのグループの何らかの一意の属性にCRUSHアルゴリズムを適用することによって配置される。

【0014】

ノードおよびストレージの障害に対する回復力を提供するために、データは、冗長的に記憶される。データまたはデータの消失訂正符号化（erasure-code）されたフラグメントのコピーが、ノードまたはストレージ・デバイスのサブセットが障害にあうときにデータがアクセス可能およびリカバリ可能なままであることができるような方法でクラスタ内の複数の位置に分散される。データのリカバリ可能な冗長な配置のためのさまざまな技術が、存在する。例示的な実施形態は、CRUSHアルゴリズムの特徴を利用し、CRUSHアルゴリズムは、クラスタ中に擬似ランダムに分散された1組の位置を計算し、同じストレージ・デバイスまたはノードへのデータの複数のコピーの配置を避けるように制約される。オブジェクト毎にオブジェクトの配置およびオブジェクトのメタデータを追跡することは、クラスタに大きな処理負荷をかける。たとえば、数百万のオブジェクトを有するストレージ・システムは、オブジェクト毎に配置を追跡するとき、過度の負荷に晒される。例示的な実施形態において、チャンクは、より大きな論理的なグループ（本

明細書においては「整合性グループ (consistency group)」または「CG」と呼ばれる)に集約され、CG内のすべてのチャンクの配置のための位置が、同じ一意の属性を用いて計算される。例示的な実施形態において、CGは、CGの「名前」にCRUSHアルゴリズムを適用することによって配置される。CGは、データを複数のオブジェクトの集合体として追跡することによって上述の問題に対処し、集約されたデータがクラスタ内に冗長的に記憶されるという付け加えられた利点をさらに提供する。その上、整合性グループは、計算の量と、データを記憶し、取り出すときにシステムが追跡しなければならないオブジェクト毎のメタデータの量とを削減する利点を有する。たとえば、メタデータは、ファイル名、ID、パーミッション、サイズ、種類、データが生成されたときにに関する情報(生成時間)、そのイベント番号の後にデータが書き込まれたイベント番号、そのイベント番号の後にデータの状態を示すために使用されるフラグが設定されるイベント番号、データのアクセス時間などのファイル属性に関する1つまたは複数の情報を含む可能性がある。

10

20

30

40

50

【0015】

一部のその他の実施形態においては、ストレージ・デバイスが、各ノードのローカルに取り付けられ、クラスタ内のその他のノードによってアクセスされ得ない。さらに、データ処理は、各ノードがローカルで利用可能であり、クラスタ内のその他ノードによってアクセスされ得ない。そのようなストレージ・システムに記憶されるデータは、ノードのうちの1つもしくは複数の障害またはノードのうちの1つのローカルに取り付けられたストレージの障害の場合にデータの可用性および回復力を提供するためにノードの間で複製される。クラスタ内のいずれかの場所に記憶されるデータは、ノード間通信接続を介してクラスタ内の他のいずれかの場所からアクセス可能であり、データの位置およびデータ・アクセス機構は、クラスタに記憶されるデータを生成し、修正し、読み取るクライアントには見えない。

【0016】

一部の実施形態において、各ノードのローカル・ストレージ・エリアは、インテリジェントなオブジェクト・ストレージ(本明細書においてはオブジェクトベース・ストレージ・デバイス(Object-based Storage Device)またはOSDと呼ばれるが、SNAおよびT10 OSD規格に合致しない)として構成される。OSDは、ディスク・ストレージと同様であるが、記憶空間のより抽象化されたビュー(view)を提供する。データの決まったサイズのブロックとして読み取りおよび書き込みを処理する代わりに、OSDは、データをオブジェクト・データ・コンテナへと編成する。各オブジェクトは、データと、オブジェクトを記述する属性情報などのメタデータとの両方を有する。

【0017】

一部の実施形態において、ノードまたはディスクの障害などの出来事がクラスタ内で起こるとき、システムは、影響を与えられたデータ・アイテムを修復しなければならない。たとえば、リカバリの場合、リカバリ・プロセスにかけられる各ノードが、そのノード自体のリカバリの管理を課せられる。リカバリは、データ修正オペレーションの信頼性の高い追跡と、ストレージ・クラスタおよびストレージ・クラスタ内のデータの状態をすべてのノードに知らせることと、リカバリオペレーションがローカルのデータに対して実行されなければならないかどうかおよびそのようなオペレーションが何を引き起こすかを独立して判定することとによって実現される。一部のその他の実施形態において、ディスクまたはノードが追加または削除されたためにクラスタの構成が変わる場合、データは、新しい記憶容量を利用するためまたは記憶容量が削減されることに適応するために再配分または再配置される必要がある可能性がある。さらにその他の実施形態においては、例えば、記憶されなければならないデータのコピーの数を増やすことによって冗長性のポリシーが変わる場合、新しいデータのコピーが生成されなければならない。そのような出来事は、本明細書においては「リカバリ」と呼ばれるカテゴリの下にグループ化される。本発明の1つの目的は、クラスタ内のさらなる障害がデータの損失をもたらす可能性があるので、

速く、効率的で、信頼性が高いリカバリを提供することであり、データの適時の効率的なリカバリを実行することによって、さらなる障害が原因であるデータの損失の可能性が、大幅に小さくされ得る。

【 0 0 1 8 】

その他の実施形態においては、データ・オブジェクトの状態が、進行中の書き込みおよび欠けているコピーであるデータ・オブジェクトへの書き込み、「良好な」ノードが「壊れている」ノードに書き込む「プッシュ」の手法の代わりに、影響を与えられるノードからの「プル」の手法を用いることによる、同期的に複製されるデータへの書き込みの結果として生じる I / O エラーからのリカバリ、グローバルな状態の依存性の連鎖を生み出すことのない、書き込み側の障害によって引き起こされた同期的に複製されるデータへの部分的な書き込みからのリカバリ、ストレージおよび / またはノードが追加または削除される場合のデータのマイグレーション、ならびにデータの冗長性のポリシーが変わるときの同期的に複製されるデータの冗長なコピーの追加および / または削除によって高い信頼性で追跡される。

【 0 0 1 9 】

さらにその他の実施形態においては、ノードの間でイベントおよび処理を調整するために、分散型の合意アルゴリズムがクラスタで実施される。合意アルゴリズムの例は、当技術分野で知られている Paxos ファミリーのアルゴリズムである。Paxos アルゴリズムにおいては、ノードのうちのいずれかが、「リーダー」として働き、システム内のあらゆるノードによって実行するためのコマンドを提案しようとする可能性がある。あらゆるそのような提案は、提案をより簡単に追跡するための対応する提案番号とともに送信される可能性がある。そのような提案番号は、ノードが実行するためにコマンドについて合意しようとして試みている特定のステップといかなる関係も持つ必要がない。最初に、リーダーが、リーダーが提出しようとする意図する提案に関する提案番号を示唆する可能性がある。そして、残りのデバイスのそれぞれは、それらのデバイスが賛成する最後の提案の指示、またはそれらのデバイスがいかなる提案にも賛成しないという指示によって提案番号のリーダーの示唆に回答する可能性がある。さまざまな応答を通じて、リーダーは、デバイスによって賛成されたいかなるその他の提案も知らない場合、以前のメッセージで示唆された提案番号を用いて、所与のクライアントのコマンドがデバイスによって実行されることを提案する可能性がある。各デバイスは、その段階で、アクションに賛成すべきかまたはそのアクションを拒否すべきかを判定する。デバイスは、そのデバイスがより大きな提案番号の別のリーダーの示唆に回答した場合にだけ、アクションを拒否すべきである。最小限の数として知られる十分な数のデバイスが提案に賛成する場合、提案されるアクションは、合意されたと言われ、各デバイスが、アクションを実行し、結果を送信することができる。そのようにして、Paxos アルゴリズムは、デバイスのそれぞれが同じ順序でアクションを実行することを可能にし、ノードのすべての間で同じ状態を維持する。

【 0 0 2 0 】

概して、Paxos アルゴリズムは、上述のように、リーダーがデバイスによって賛否の意思表示をされた前の提案を知ることができる最初のフェーズと、リーダーが実行するためのクライアントのコマンドを提案することができる第 2 のフェーズとを有する 2 つのフェーズにあると考えられる可能性がある。リーダーは、前の提案を知ると、第 1 のフェーズを繰り返す必要がない。その代わりに、リーダーは、第 2 のフェーズを継続して繰り返し、複数のステップで分散型コンピューティング・システムによって実行され得る一連のクライアントのコマンドを提案する可能性がある。そのようにして、各ステップに関して分散型コンピューティング・システムによって実行されるそれぞれのクライアントのコマンドは、Paxos アルゴリズムの 1 つのインスタンスと考えられ得るが、リーダーは、次のステップに関する別のクライアントのコマンドを提案する前に所与のステップのための提案されたクライアントのコマンドに対してデバイスが賛否の意思表示をするのを待つ必要がない。

【 0 0 2 1 】

実施形態の一部においては、クラスタ内の各ノードが適切な順番でデータを処理していることと、各ノードが同じ状態を維持することとを保証するために Paxos アルゴリズムが使用される。このようにして、Paxos アルゴリズムに従う個々のノードは障害にあう可能性があるが、障害の間にクラスタ内でいかなるオペレーションも失われない。Paxos アルゴリズムを実施する実施形態においては、個々のノードのいずれかが、リーダーとして働き、クラスタ内のあらゆるその他のノードに実行するためのオペレーションを提案しようとする可能性がある。あらゆるノードが同じ順序で同じコマンドを実行することを必要とすることによって、Paxos アルゴリズムは、クラスタのノードの間の同期化を実現する。加えて、クラスタ内のノードが障害にあうかまたはそうではなくクラッシュする場合、クラスタは、Paxos のメッセージを用いて知らされる。参照により本明細書に援用される、Leslie Lamport による「Paxos Made Simple」、ACM SIGACT News (Distributed Computing Column) 32、4 (通巻号数 121、2001 年 12 月) 51~58 を参照されたい。参照により本明細書に援用される、Tushar Deepak Chandra、Robert Griesemer、および Joshua Redstone による「Paxos Made Live - An Engineering Perspective (2006 Invited Talk)」、Proceedings of the 26th Annual ACM Symposium on Principles of Distributed Computing、ACM press (2007) も参照されたい。

10

20

30

40

50

【図面の簡単な説明】

【0022】

【図 1】本発明の一実施形態によるクラスタ化されたシェアード・ナッシング・ストレージ・システムの図である。

【図 2】本発明の一実施形態によるノードのブロック図である。

【図 3】本発明の一実施形態による整合性グループと、チャンクと、スライスとの間の関係を示すブロック図である。

【図 4】本発明の一実施形態によるシェアード・ナッシング・ストレージ・システム上のデータ・オブジェクトの配置を決定するためのストレージ構成マップの例を示す図である。

【図 5】本発明の一実施形態による CRUSH によって生成されるディスク・リストの例を示す図である。

【図 6】本発明の一実施形態によるクラスタのすべてのノードの状態情報の例を示す図である。

【図 7】本発明の一実施形態によるマップ変更プロセスに関する高レベルの流れ図である。

【図 8】本発明の一実施形態による各ノードが記憶するイベント・リストの例を示す図である。

【図 9】本発明の一実施形態による図 7 のマップ変更プロセスのフェーズ 1 のプロセスの流れ図である。

【図 10】本発明の一実施形態によるフェーズ 1 中にマイグレーション要求を受信するノードのプロセスの流れ図である。

【図 11】本発明の一実施形態による図 7 のマップ変更プロセスのフェーズ 2 のプロセスの流れ図である。

【図 12】本発明の一実施形態によるフラグ・テーブルの例を示す図である。

【図 13】本発明の一実施形態によるチャンク・スタブ (stub) 管理情報の例を示す図である。

【図 14】本発明の一実施形態によるノードのメモリに記憶されるビットマップの例を示す図である。

【図 15】本発明の一実施形態によるマイグレーション中の WRITE 要求のプロセスの

流れ図である。

【図 16】本発明の一実施形態によるマイグレーション中の W R I T E 要求のプロセスの流れ図である。

【図 17】本発明の一実施形態による、通常の条件の間に I / O オペレーションを受け取るノードのプロセスの流れ図である。

【図 18】本発明の一実施形態による、マイグレーション中の I / O オペレーションが完了した後にデータ・チャンクをクリーンに設定するプロセスの流れ図である。

【図 19】本発明の一実施形態による、障害の後、オンラインに復帰するときのターゲット・ノードのプロセスの流れ図である。

【図 20】本発明の一実施形態によるマイグレーション中のソース・ノードの障害のリカバリのプロセスの流れ図である。

【図 21】本発明の一実施形態による図 7 のマップ変更プロセスのフェーズ 3 のプロセスの流れ図である。

【図 22】本発明の一実施形態による整合性グループに関するポリシー・テーブルを示す図である。

【図 23】本発明の一実施形態による冗長性のレベルのポリシーの変更のプロセスの流れ図である。

【図 24】本発明の一実施形態による、クラスタに再び加わった後の部分的に書き込まれたデータのノードのリカバリのプロセスの流れ図である。

【図 25】本発明の一実施形態による、クラスタに再び加わった後の部分的に書き込まれたデータのノードのリカバリのプロセスのメッセージ・シーケンスを示す図である。

【図 26】本発明の一実施形態による (I / O) オペレーションの受け取りを扱うためのホストとノードとの間の通信フローを示す図である。

【図 27】本発明の一実施形態による I / O オペレーションを扱うためのホストとノードとの間の代替的な通信フローを示す図である。

【発明を実施するための形態】

【0023】

本明細書において検討される実施形態は、本発明の 1 つまたは複数の例を示す。本発明のこれらの実施形態が図を参照して説明される説明されるとき、説明される方法および / または特定の構造のさまざまな修正または適応が、当業者に明らかになる可能性がある。本発明の教示に依拠し、これらの教示が当技術分野を進歩させたすべてのそのような修正、適応、または変形は、本発明の範囲内にあると考えられる。したがって、本発明は本明細書において示される実施形態のみにまったく限定されないことが理解されるので、この説明および図面は限定的な意味にとられるべきでない。別途示されない限り、同様の部分および方法のステップは、同様の参照番号によって参照される。

【0024】

導入

本発明は、ネットワーク接続ストレージ (N A S) のアーキテクチャの基盤として実装される可能性があり、分散型データ・プラットフォームで使用される可能性がある。概して、図 1 に示されるように、ネットワーク 30 に接続されたノード 10 上のファイル・システムにおいてデータ・オブジェクトを分散させ、記憶するシェアード・ナッシング分散型ストレージ・システムが、提供される。

【0025】

図 1 に示されるように、クラスタの各ノードは、ネットワーク 30 を通じて接続される。1 つまたは複数のクライアントまたはホスト 20 が、シェアード・ナッシング・ストレージ・システムに結合される。ホスト 20 は、概してコンピュータであるが、ネットワークを介して接続された別のストレージ・システムまたはサーバである可能性もある。ホスト 20 は、たとえば、I / O オペレーションをストレージ・システムに送る。I / O オペレーションは、例えば、書き込みオペレーション (W R I T E)、読み取りオペレーション (R E A D)、削除オペレーション (R E M O V E)、またはトランケート (T R U C

A T E) オペレーションである。チャンク・データを修正する I / O オペレーション (そのようなオペレーションは W R I T E である) は、そのチャンクを記憶するノードにグローバル一意キーを供給する。図 2 に示されるように、各ノード 1 0 は、そのノード 1 0 が記憶する各チャンクに関する I / O キーを保有する I / O キー・リスト 9 5 をメモリ 5 0 に保有する。加えて、I / O オペレーションがダウンしているノード 1 0 に対して要求される場合、ノード 1 0 は、クラスタ内のその他のノードにそれらのその他のノードの最新の I / O キー・リストに関して問い合わせることによって、オンラインに復帰するときに I / O オペレーションを受け取る。アクティブなノード 1 0 またはディスク 6 5 は、本明細書においてオンラインと呼ばれ、一方、非アクティブなノード 1 0 またはディスクは、本明細書においてオフラインと呼ばれる。ノード 1 0 のいずれかが、必要に応じて手動でオフラインもしくはオンラインにされる可能性があり、またはそれらのノードは、ソフトウェアもしくはハードウェアの障害によってオフラインになる可能性がある。

10

【 0 0 2 6 】

典型的には、ホスト 2 0 は、パーソナル・コンピュータ (P C)、ワークステーション、ラップトップ、携帯情報端末 (P D A)、サーバ、メインフレームなどのコンピュータ・システムである。ホスト 2 0 は、N F S、C I F S、H T T P、F T P などのファイル・アクセス・プロトコルを用いて遠隔のファイルおよびファイル・システムにアクセスするように構成される。

【 0 0 2 7 】

ノード 1 0 は、P C、ワークステーション、サーバ、メインフレームなどである可能性があり、またはディスク・アレイもしくはストレージ・アプライアンスなどの、本明細書において説明されるシステムを実装するのに十分な組み込まれたコンピューティング能力を有するストレージ・デバイスである可能性がある。ノードは、ローカル・ファイル・システム、ネットワーク接続ストレージ (N A S)、ストレージ・エリア・ネットワーク (S A N)、データベースなどの上のファイル・システム内のファイルに関連する情報を記憶し得る。さらに、ノードは、ファイル・システムにファイルを記憶するために任意のハードウェアおよび/またはソフトウェア要素から構成される可能性があり、N T F S、E X T、X F S、G F S などの、ファイルを記憶するための 1 つまたは複数のファイル・システムを実装する可能性がある。

20

【 0 0 2 8 】

クラスタにおいて、ファイル・システム自体は、特定のノードに物理的に結びつけられないので、1 つもしくは複数のノードまたはクラスタ全体に広がる可能性がある。データは、ノード 1 0 に記憶され、ノード間通信接続を介してクラスタ内のその他のノード 1 0 によってアクセスされ得る。ネットワーク・インターフェース 4 0 および基礎を成すネットワーク通信の技術的な態様は、本発明の範囲内になく、したがって、それらの態様は、詳細に説明されない。クラスタ内のノード 1 0 の間のデータのマイグレーション中、ホスト 2 0 からの I / O 処理は、妨げられない。

30

【 0 0 2 9 】

ストレージ・システム内のノード 1 0 に記憶されたデータは、ノード 1 0 のうちの 1 つもしくは複数の障害またはノード 1 0 に取り付けられたディスク 6 5 の障害の場合にデータの可用性および回復力を提供するためにノード 1 0 の間で複製される。すべてのノード 1 0 は、その他のノード 1 0 のステータスと、そのノードのディスク 6 5 およびその他のノード 1 0 のディスク 6 5 上のデータの状態とを状態情報 1 0 0 を介して知らされる。消失訂正符号化方法が、ストレージ・システムにおいて使用される可能性がある。たとえば、リード・ソロモン方式においては、データのブロックに対する計算が、消失訂正符号化されたブロックを生成する。数多くのその他の好適な冗長なまたは消失訂正符号化されたストレージ方式が、当業者に明らかであろう。各ノード 1 0 は、処理能力と記憶能力との両方を有する可能性がある。クラスタ内のノード 1 0 は、ローカル・ファイル・システム、N A S、S A N、データベースなどの上のファイル・システム内のファイルに関連する情報を記憶し得る。クラスタに記憶されたファイルを提供するための容量および帯域幅を

40

50

スケーリングするために、ノード 10 が追加されるかまたはクラスタから削除される可能性があり、ノード 10 のディスク 65 が追加または削除される可能性がある。

【0030】

図 4 のように、ノード 10 のクラスタにわたるストレージ・システム上のデータの分散された記憶に関するレイアウトを決定するために、ストレージ構成マップ 70、71 が、すべてのノード 10 によって利用される。たとえば、ノード 10 またはディスク 65 が追加されるとき、クラスタにおけるノードまたはディスクの追加に基づく新しいマップ 71 をコミットするために、マップ変更提案プロセスが開示される。古いおよび新しいマップ 70、71 に基づいて、各ノード 10 は、データ・オブジェクトの配置または位置を知る。古いマップ 70 と新しいマップ 71 との両方は、ノード 10 における I/O 処理およびレイアウトの参照のために利用され得る。

10

【0031】

図 4 に示されるマップ 70、71 は、整合性グループ ID (CGID 77) に基づいて CRUSH アルゴリズムがデータのチャンクの配置を決定するノード 10 およびディスク 65 またはボリュームの集合である。ストレージ・システムは、図 2 に示された論理デバイス (LDEV) 66 を用いて構成された 1 つまたは複数の論理ストレージ・エリアまたは論理ボリュームを提供する。ストレージ・デバイス 67 は、論理ボリュームと LDEV 66 との間の対応 (関係) を管理する。ディスク 65 または論理ドライブ上のボリュームが、マップ 70、71 上の CG のチャンクに関する位置として提供され得る。したがって、マップ 70、71 における所与の CG のチャンクの配置または位置は、ディスク 65 の代わりにボリュームまたは論理ボリュームを用いて定義される可能性もある。

20

【0032】

各 CG は、CG を特定する一意識別子、CGID 77 を有する。CG が生成されると、生成するノード 10 が、生成された CG に関する CGID 77 を生成する。マップ 70、71 は、データ・ストレージ・システム上のデータ・オブジェクトの分散された記憶に関するレイアウトを決定する際に利用される。CRUSH は、ノード 10 およびディスク 65 の CG およびその CG がマップ 70、71 に基づいてクラスタ内に含むストレージ・オブジェクトの配置を選択する。CG は、CGID 77 を共有し、したがって、ノード 10 およびディスク 65 上の同じ配置を有するデータ・チャンクからなるファイルなどのデータ・オブジェクトのグループである。データ・オブジェクトは、その他のデータ・オブジェクトと別々に識別可能であり、ストレージ・システムに転送可能である有限長のデータであり、テキスト・ファイル、画像ファイル、またはプログラム・ファイルなどである可能性がある。CG およびそれらの CG の対応する複製が、クラスタ内の複数のノード 10 にわたって記憶される。図 3 は、CG と、チャンクと、スライスとの間の関係の図を示す。図 3 に示されるように、CG は、複数のチャンクから成り、各チャンクは、限られた最大サイズ (たとえば、64 MB) のデータ領域である。チャンクは、信頼性のためにノード 10 中に複数複製されて記憶されるかまたは消失訂正符号化される。各チャンクは、スライスと呼ばれるデータの領域にさらに分割され得る。各スライスは、決まったサイズである可能性がある。

30

【0033】

CRUSH アルゴリズムは、マップ 70、71 からデータが記憶されるディスク 65 のリストを生成するために使用され得る。任意のノード 10 が、CRUSH アルゴリズムを実行することができる。ノード 10 は、CGID 77 およびいくつかの必要とされる一意の配置位置を含む情報を CRUSH アルゴリズムに渡す。CRUSH は、例えば、ディスク 65 の明示的に順序付けられたリストによって応答する。CRUSH は、同じマップ 70、71 内の同じ CGID 77 に関していつもディスク 65 の同じシーケンスが生成されることを保証する。つまり、CRUSH は、3 つのディスク {A, B, C} またはディスク・リスト 76 {A, B, C, D} を生じる 4 つのディスクのディスク・リスト 76 を返す可能性がある。リスト 76 のディスク 65 は、アクセス可能である可能性があり、またはソフトウェアもしくはハードウェアの問題、たとえば、特定のディスク 65 が物理的に

40

50

取り付けられているノード 10 のクラッシュが原因でアクセス不可能である可能性がある。クラッシュによって生成されるリスト 76 の第 1 の使用可能なディスク 65 は、「ファースト・アップ (first up)」ディスク (ボリューム) と呼ばれる可能性がある。

【0034】

図 4 は、古いマップ 70 および新しいマップ 71 の例を示す。古いマップ 70 は、以前のマップ 70 のバージョンに関する CRUSH マップを含む 1 組のノード 10 およびディスク 65 (ボリューム) である。一方、新しいマップ 71 は、提案される新しいマップ 71 のバージョンに関する CRUSH マップを含む 1 組のノード 10 およびボリューム 65 として定義され得る。古いマップ 70 は、新しいマップ 71 がノード 10 によってコミットされるまでデータ・オブジェクトのレイアウト構成として使用される。言い換えると、マップ 70 は、クラスタ内のデータの配置のために使用される現在のマップに対応し、一方、マップ 71 は、クラスタ内のデータの新しい提案される配置である。図 4 のキーで示されるように、ノード 10 は、ノード識別子、たとえば、番号 1、2、3、4 を用いて示され、ディスク 65 は、大文字 A、B、C、D、E を用いて示される。各ディスク 65 は、対 (ノードの番号, ディスクの文字) によって特定される。たとえば、1A は、ノード 1 のディスク A を表す。

【0035】

図 5 は、CGID77「1」に関する古いマップ 70 および新しいマップ 71 を用いて CRUSH によって生成された例示的なディスク・リスト 76 を示す。図 4 および 5 に示されるように、CGID77「1」に関する古いマップ 70 の下のディスク・リスト 76 は、1A、3C、および 4F を含む。一方、新しいマップ 71 の下では、CGID77「1」は、ディスク 2B、3C、および 5A に記憶される。提案される (新しい) マップ 71 の下で、CGID77「1」に関する CRUSH によって生成されたディスク・リスト 76 は、2B、3C、および 5A を含む。

【0036】

ノード 10 は、ファイルシステム、ネットワーク接続ストレージ (NAS)、ストレージ・エリア・ネットワーク (SAN)、データベースなどの内のデータに関連するメタデータを記憶する。メタデータは、ファイル名、ID、パーミッション、サイズ、種類、データが生成されたときに関する情報 (生成時間)、そのイベント番号の後にデータが書き込まれたイベント番号、そのイベント番号の後にデータの状態を示すために使用されるフラグが設定されるイベント番号、データのアクセス時間などのファイル属性を含む。図 2 に示されるように、各ノード 10 は、オブジェクト・ストレージ・デバイス (OSD) 60、ネットワーク・インターフェース 40、プロセッサ 45、(本明細書においては SA と呼ばれる) 管理ユーティリティ 55、メモリ 50、および複数のディスク 65 を有する。ネットワーク・インターフェース 40 は、ネットワーク上のホストおよびその他のノードと通信し、ホスト 20 から I/O 処理を受信し、その I/O 処理に応答する役割を担う。加えて、Paxos および CRUSH アルゴリズム、マップ 70、71、メタデータ、ならびにビットマップ 75 が、メモリ 50 において構成され、管理される。

【0037】

各ノードのメモリは、CRUSH および Paxos アルゴリズムを実行し、各 CG およびその CG のチャンクのメタデータに関連する情報、フラグ・テーブル 120、I/O キー・リスト 95、状態情報 100、ビットマップ 75、チャンク・スタブ管理情報 105、イベント・リスト 125、ならびにチャンク・コピー・リスト 115 を記憶し、管理するために使用される。

【0038】

各ノード 10 は、独立したストレージ・デバイス 67 を有する。ストレージ・デバイス 67 は、複数のストレージ・ドライブまたはディスク 65 を含む。各ディスク 65 は、たとえば、SAS (シリアル・アタッチド SCSI)、SATA (シリアル ATA)、FC (ファイバ・チャネル)、PATA (パラレル ATA)、および SCSI などの種類のハ

10

20

30

40

50

ード・ディスク・ドライブ、半導体ストレージ・デバイス（SSD）など、またはファイバ・チャネル・ネットワーク上のストレージ・エリア・ネットワーク（SAN）を介して遠隔でアクセスされるSCSI論理ユニット（LUN）もしくはTCP/IPネットワーク上でアクセスされるiSCSI LUNとして与えられるストレージ・デバイス、またはその他の永続的なストレージ・デバイスおよびメカニズムである。

【0039】

各ノード10は、I/Oを実行するためにホスト20が通信するレイヤであるストレージ・アクセス・レイヤ（SAS）を提供する管理ユーティリティを有する。SASは、ストレージのミラーリングおよび整合性を扱う。SASは、分散型I/Oシステムの「イニシエータ」部分である。加えて、SASは、各ノード10のOSDからのデータを要求することができる。各ノード10のOSD60は、ノード10上のローカルI/Oを扱い、ストレージ・デバイス67に対するオペレーションを実行する。OSD60は、たとえば、ストレージ・デバイス67からデータを読み取るかまたはストレージ・デバイス67にデータを書き込むときにストレージ・デバイス67との通信を実行する。OSD60は、CGに関して、そのOSD60がCGのチャンクをどのディスク65に記憶するかを決定することができる。マップ70、71が、各ノード10のSASにおいて生成される。また、あらゆるノード10は、マイグレーションが現在実行されているかどうかを示すマップ・マイグレーション状態を知っている。CGは、そのCG自体がリカバリされている途中であるかどうか、またはCGがマイグレーション状態であるかどうかをシステムに知らせる状態指示によって印付けられる。加えて、SASは、マップ70、71を管理する。SASのレイヤは、マップ70、71が古いマップ70であるかまたは新しいマップ（提案されるマップ）71であるかに関わらず、マップ70、71を使用する。

【0040】

各ノード10は、その他のノード10によって受信され、ノード10のクラスタの個々のビューを構築し、維持するために使用される状態情報100をネットワーク上で提供する。図6は、クラスタ10のノードがアクティブであるかどうかおよびノードのディスクがデータの記憶のために利用可能であるかどうかについての情報を含む、ノードおよびノードのディスクの状態（ステータス）情報100の例である。加えて、状態情報100は、それぞれのローカルに取り付けられたストレージ・デバイス67内の個々のディスク65の可用性を含む。当業者は、状態情報100が異なる属性を含む可能性があり、または更新ステータスもしくはエラー・ステータスなど、異なるように定義される可能性があることを認識する。図6に示される例においては、ノード1がオンラインであり、ディスクA～Cがオンラインであり、一方、ノード2のディスクB、D、およびEはオフラインである。状態情報100は、動的であり、各ノードのメモリに記憶される。

【0041】

ノード10のディスク65または新しいノード10などの追加的な記憶容量がクラスタに追加されるとき、ディスク65の削除など、記憶容量が削減されるか、もしくはノード10が削除される（所与のノード10に取り付けられたすべてのディスク65をオフラインにする）とき、または重い負荷がかかっているノード10もしくはディスク65から十分に利用されていないノードに割り当てを振り向けるようにディスクの重み係数が変更されるとき、マップ変更が行われなければならない、したがって、クラスタに関するI/O処理がクラスタ構成に対する変更を反映するために新しいマップ71に従って実行される。マップ変更は、一部のCGが新しい位置（ノード10/ディスク65）にマッピングされるようにし、それらのCG（すなわち、CGのすべてのチャンク）のコピーが新しいマップ71がコミットされる前にそれらのCGの新しい位置に再配置されることを要求する。マップ変更の下でのマイグレーション中、ホスト20からのI/O処理は、妨げられない。

【0042】

マップ変更プロセス

図1のストレージ・システムの分散型シェアード・ナッシング・アーキテクチャにおい

10

20

30

40

50

て、各ノード10は、集中型の調整または監督なしにそのノード10自体のデータのリカバリ・プロセスを実行する役割を担う。さらに、Paxosを用いるマップ変更イベント配信アルゴリズムを実施することによって、各リカバリ・イベントが、各ノード10に順に配信され、各ノード10が、起こったすべてのイベントを受信する。

【0043】

図8は、各ノード10がメモリ50において管理するイベント・リスト125の例を示す。各ノード10のSA55がPaxosによってイベントを受信するとき、そのイベントは、時間順にイベント・リスト125に記憶される。Paxosは、イベントが処理されるために届けられる順序を管理する。ノード10は、オンラインに復帰するとき、オフラインであった間に逃した可能性があるイベントおよびタスクのリストを取得するためにクラスタ内のそのノード10のピアに問い合わせる。ノードは、イベントを受信し、それらのイベントをイベント・リスト125に記憶する。図8に示されるように、ノード1は、MAP__CHANGE__PROPOSALイベントを受信し、続いてMAP__MIGRATEイベントを受信し、その後の時間にMAP__COMMITイベントを受信した。したがって、ノード10は、順番にイベントに対応するプロセスを行わなければならない。ノード10は、現在のフェーズが完了するまで図7に示されるように後のフェーズに移ることができない。図7に示される3つのフェーズのそれぞれが、以下の通り詳細に説明される。

【0044】

上述のように、たとえば、クラスタの記憶容量がクラスタへのノードの追加もしくは削除またはノードへのストレージ・デバイスの追加もしくは削除によって変更されるとき、CRUSHアルゴリズムが、クラスタ中のデータの配置に関する新しいマップを決定するために使用される。既存のマップ（たとえば、古いマップ）から提案されるマップ（たとえば、新しいマップ）への移行は、マップ変更プロセスと呼ばれる。マップ変更プロセス中、データが、新しいマップに合致するために1つまたは複数のノードに再配置される可能性がある。データが再配置される間、データが引き続きクラスタによってホスト20に利用され得るようにされる、I/O処理が移行中に妨げられないことが重要である。マップ変更プロセスは、3つのフェーズで行われる。

【0045】

図7は、マップ変更プロセスの抽象的な図を示す。概して、フェーズ1 100は、マップ変更の提案からなる。マップ変更の提案が行われると、各ノード10は、マップ変更の準備をし、提案に応答する。マップ変更は、ノードによって提案されるか、管理命令によってトリガされる。マップ変更の提案は、クラスタ内のオンラインであるすべてのノード10にPaxosに従って配信される。各ノード10は、あらゆる継続中のプロセスを終わらせることなどによってそのノード10自体をマップ変更のために準備する第1のステップを経なければならない。各ノード10は、そのノード10自体がフェーズ1に関するそのノード10のそれぞれの処理を済ませたと考えるときに、マップ変更の提案を受け付けたことをPaxosに従って報告する。マップ変更の提案元は、すべての「ダウンしていない(up)」ノード（たとえば、オンラインであり、応答可能であるノード10）がマップ変更の提案に応答し、マップ変更が受け付け可能であることを提案元ノードに示すまでフェーズ2に進むことができない。いずれかのノード10が提案を拒絶する場合、提案は失敗する。提案がクラスタ内のすべてのダウンしていないノードによって受け付けられないことに応じて、マップ変更プロセスのフェーズ2 200に入る前に別のマップの提案が生成され、受け付けられる必要がある。

【0046】

別の実装においては、最大の許容される時間よりも長くダウンしているクラスタのノード10が、構成から削除される可能性がある。時間の最大の量は、時間の予め決められた量である可能性がある。ノード10の削除は、マップ変更をトリガし得る。この場合、削除は、マップ変更の提案元によりマップ変更を失敗させる理由であると解釈され得る。マップ変更プロセスのフェーズ2 200において、マップ変更が行われるために必要とさ

10

20

30

40

50

れる処理を勧めるようにすべてのノード10に命令するメッセージが、Paxosに従って送信される。つまり、フェーズ2は、クラスタ内でチャンクのマイグレーションおよび書き込みの代理処理(proxying)をトリガする。フェーズ2は、すべてのダウンしていないノード10がそれらのノード10が新しいマップを実施するために必要とされるそれらのノード10のそれぞれのマイグレーション処理を完了したという報告をPaxosに従って返すときに終了する。

【0047】

フェーズ3 300は、新しいマップ71が実施されていることをすべてのノード10に知らせることによってマップ変更を「コミットすること」からなる。これは、すべてのノード10のSA55に古いマップ70を使用することを止めさせ、新しいマップ71を使用し始めさせる。第3のフェーズにおいて、各ノード10に関するOSD60は、新しいマップ71の下でもはや記憶されない古いマップの下で記憶されたデータを削除する。

10

【0048】

フェーズ1

フェーズ1が、図9(フェーズ1)に示される流れ図を参照して説明される。フェーズ1の初めに、あらゆるノード10が、マップ変更の提案元ノードからPaxosに従ってMAP__CHANGE__PROPOSALイベント101を受信する。これが、フェーズ1のメッセージである。図9において、提案されるマップは、説明を目的として、新しいマップ71を指す。各ノード10によってイベント101が受信されると、新しいマップが、ノードのメモリ50に記憶されるあらゆるノード10において生成される(ステップ105)。

20

【0049】

ステップ110において、新しいマップ71に基づいて、各ノード10が、新しいマップ71へのマップ変更によって影響を与えられるいずれかのCGをそのノードが記憶するかどうかを判定する。判定は、各CGに関して行われる。ノードによって記憶されるCGは、古いマップ70の下での位置と比較されるときに新しいマップ71の下でのCGに関する位置の変化がある場合に、新しいマップ71によって影響を与えられる。古いマップ70と新しいマップ71との間の比較は、各CGに関するCGID77を調べ、CGID77に対応する(古いマップ70および新しいマップ71に関する)ディスク・リスト76を比較することによって行われる。各ノード10は、新しいマップ71によって影響を与えられる各CGに関して図9のステップに従う。ノード10は、新しいマップ71の下で影響を与えられるCGを持っていないか、マップ変更によって影響を与えられるCGを持っているかのどちらかである。ノード10が影響を与えられるいかなるCGも持たない場合、ノードは、ステップ115において、Paxosに従ってそのノードの受け付けをポストする、つまり、マップ変更の提案を受け付けた。マップ変更がノード上で今やCGに影響を与える場合、処理は、ステップ140に続く。

30

【0050】

ノード10は、古いマップ70でノード10のローカル・ボリューム65にコピーを記憶するが、新しいマップ71の下ではそのコピーをローカルに記憶しない場合、新しいマップの提案のそのノード10の受け付けをポストし、新しいマップ71がコミットされるまでI/Oオペレーションを受け付け続ける。ノード10は、古いマップ70と新しいマップ71との両方下でそのノード10のローカル・ボリューム65にCGのコピーを有する場合、新しいマップ71のそのノード10の受け付けをポストする。また、ノード10は、(提案される)新しいマップ71が有効にされるまでI/O要求を受け付け続け、これは、後でフェーズ2および3において説明される。

40

【0051】

フェーズ1中に、ノード10は、そのノード10が古いマッピング・シーケンス70の第1の「ダウンしていない」ボリュームを含むかどうかをやはり判定する。ダウンしていないボリュームは、オンラインであり、応答可能であり、データを新しい位置にマイグレーションすることができるノード10のボリューム65と呼ばれる。ダウンしていないボ

50

リュームは、そこから新しい位置にデータがマイグレーションされるボリュームである可能性がある。ダウンしていない複数もボリュームが存在する可能性があるが、1つが、第1のダウンしていないボリュームとして指定される。ダウンしていないボリュームは、CGID77に基づいてCGに関してCRUSHによって生成されるディスク・リスト76の第1のディスク65としてノードによって決定され得る。

【0052】

ステップ140において、ノード10が第1のダウンしていないボリュームを有するとノード10が判定する場合、そのノード10は、新しいマップ71の下の整合性グループのその他のメンバーに、それらのその他のメンバーが新しいマップ71の下で記憶されるCGを有することを知らせる役割を担うノード10である。第1のダウンしていないボリュームを有するノードは、新しいマップ71を用いてCGのチャンクに関する新しい位置を特定する。ノード10は、ステップ145において、新しいマップの下でのCGのメンバーであるノードに、それらのノードが第1のダウンしていないボリュームのノードからCGを受信すべきであることを示すCG__MIGRATION__REQUESTメッセージを送信する。フェーズ1の下では、ステップ145において第1のダウンしていないボリュームが新しいメンバーのノード10に知らせない限り、新しいマップの下でのCGを記憶すべきである新しいマップ71の下のノードのメンバーが、それらのノードのメンバーがCGのチャンクに関する新しい位置であることを知るその他の方法が存在しない。CG__MIGRATION__REQUESTは、マイグレーションされるべきであるCGID77のリストからなる。加えて、チャンクに関するフラグ・テーブル120を含むCGに対応するメタデータが、CG__MIGRATION__REQUESTとともに送信され得る。CG__MIGRATION__REQUESTを受信するノードのOSD60は、ノードが新しいマップの下での複数のCGを記憶すべきである場合、1つまたは複数のノード10の複数のダウンしていないボリューム65から複数のそのような要求を受信し得る。OSD60は、ノード10がCGを記憶する位置のリストを生成すれば十分なので、どのノード10が各要求を送信したかを追跡する必要はない。

【0053】

図10は、図9のステップ145においてCG__MIGRATION__REQUESTを受信するノード10に関する処理を示す流れ図を示す。ステップ160においてCG__MIGRATION__REQUESTを受信するノードのOSD60は、ステップ161において、CGが既に存在し、アクティブであるかどうかを判定する。CGが存在するかどうかを判定する際、ノードは、CGが実際に存在するかどうか、またはBEING__CONSTRUCTEDとして印を付けられたCGがOSD60のストレージ・デバイス内に既に存在するかどうかを判定する。CGは、そのCGに関連するエラーを持たず、ホストに書き込むことおよび/またはホストから読み取ることが現在可能である場合にアクティブである。ステップ161において、既に生成されたCGが存在しない、および/またはそのCGがエラーを有する場合、ノード10は、ステップ165に進み、ステップ165は、新しいマップの下で記憶される各CGに関するスタブおよびCGの経路の階層を生成することである。ステップ161における判定がはいである場合、このCGIDに関するCGのスタブがノードに既に存在するためノード10がそのCGのスタブを生成する必要が確かにあるので、プロセスはステップ175に進む。

【0054】

また、CGの階層は、そのCGの階層が記憶しなければならないCGに関するCGID77を含む。CGのメタデータが生成され、ステップ170においてOSD60によってBEING__CONSTRUCTEDとして永続的に印を付けられる。CGのスタブの階層が生成されると、ノードは、ステップ175において、そのノードにCG__MIGRATION__REQUESTを送信したノードにACKメッセージを送信することによってマイグレーション要求に肯定応答する。ACKメッセージは、第1のダウンしていないボリュームのノードによって受信され(ステップ150)、ノードは、そのノードがマイグレーション要求を送信したすべてのノード10からACKメッセージを受信すると、ステ

マップ155において、Paxosに従って提案に対する受け付けをポストすることによってマップ変更提案イベントに応答することができる。フェーズ1は、クラスタ内のすべてのノード10がマップ変更の提案を受け付けたときに正常に完了する。たとえば、いずれかのノード10がエラーのために提案を受け付けることができない場合、提案は失敗し、提案元はMAP__CHANGE__CANCELイベントを送出しなければならない。

【0055】

マップ変更の取り消し

提案が失敗する場合、あらゆるノード10が、Paxosに従ってMAP__CHANGE__CANCELイベントを受信する。MAP__CHANGE__CANCELイベントが起こると、マップ・マイグレーション状態がクリアされ、各OSD60が提案された新しいマップ71に関して構築されたCGを有するかどうかをそのOSD60が調べる。ノード10がMAP__CHANGE__CANCELイベントを受信すると、構築された新しいマップ71が無効化される。さらに、新しいマップ71の下でCGを受信することを予想するCGの階層を生成したノード10が、CGの階層を削除する。MAP__CHANGE__CANCELイベントがすべてのノード10に送信され、ノード10がクラッシュしたかまたはその他の理由でオフラインであるためにいずれかのノード10がイベントを受信しない場合、イベントを受信しなかったノード10は、クラスタのオンラインのノードによって保有されるイベント・リスト125に従って、そのノード10がオンラインに復帰するときにメッセージを受信する。加えて、マップ変更が取り消されると、SA55は、新しいマップ71のメンバーにI/Oオペレーションの要求を送信することを止める。

【0056】

フェーズ2

フェーズ2は、Paxosに従ってクラスタのノード10に送信されるマップ・マイグレーション指令で始まる。図11のステップ201において、あらゆるノード10が、フェーズ2のメッセージであるMAP__MIGRATEイベント・メッセージを受信する。ステップ202において、ノード10が、マップ・マイグレーション状態を「1」に設定する。マップ・マイグレーション状態は、各ノード10に関するOSD60およびSA55がマップ・マイグレーションが開始されたことを知るように設定される。

【0057】

ステップ205において、各ノードのOSD60が、それらのOSD60がBEING__CONSTRUCTEDとして印を付けられたいずれかのCGを有するかどうかを判定する。ノード10がそのようなCGを持たない場合、ステップ210において、ノードが、Paxosに従ってACKメッセージを直ちにポストしてマップ・マイグレーション・イベントに肯定応答する。BEING__CONSTRUCTEDとして印を付けられたCGを有するノード10のOSDは、古いマップ70の下でボリュームに記憶されたチャンクを有するダウンしていないボリューム(ノード10)からBEING__CONSTRUCTEDとして印を付けられたCGに属するチャンクをコピーすることによってコピー・プロセスを開始する。ステップ220において、受信ノードに関するOSD60が、フェーズ1においてCG__MIGRATION__REQUESTとともに送信されたCGID77のリストを参照することによってBEING__CONSTRUCTEDとして印を付けられたCGに関するCGを有する古いマップ70のノードのメンバーからのチャンクのリストを要求する。古いマップ70のメンバー(たとえば、記憶されたチャンクを有するノード)は、チャンクのリスト(すなわち、コピー・リスト)115を新しいマップのメンバーに送信する。

【0058】

代替的に、それぞれの新しいマップのノード10(要求を受信するノード)が、基礎を成すファイル・システムのキャッシュの効果を利用するために、同じ古いマップのノード10がチャンクのリスト115および各チャンクに関するスライスを送信することを要求する。ある場合、チャンク(スライス)を受信すべきであるそれぞれの新しいマップ70のノードが、古いマップの下でどのダウンしていないボリュームからそのノードがチャン

クを受信するかを独立に判断する。異なるダウンしていないボリュームが、同じチャンクを受信すべきであるそれぞれの新しいマップのノードによって選択される可能性がある。別の場合、チャンクを受信すべきであるそれぞれの新しいマップのノードが、チャンクを受信するために古いマップの下と同じダウンしていないボリュームのノードを選択する。後者の場合、新しいマップのノードに送信されるチャンク（スライス）が（以前のディスクの読み取りにより）選択されたダウンしていないボリュームのノードのメモリにあることになるためにキャッシュの効果が利用される可能性があり、したがって、転送されるデータが以前の転送によりノードのメモリに既にある場合、ディスクの読み取りが必要とされない。例として、古いマップ70のノード1および3が新しいマップ71のノード2および5に転送される同じデータを記憶する場合、ノード2と5との両方が、データを転送するためにノード1にのみRECOVER_Y__READ要求を送信する。ノード2がノード1からのデータを要求し、それから、ノード5がノード1からの同じデータを要求する場合、ノード1は、データがノード1のメモリにあることになるので、ノード5にデータを送信するためにディスクの読み取りを実行する必要がない。

【0059】

ステップ225において、受信ノード10が、コピー・リスト115を受信し、OSD60が、メモリ50内にコピー・リスト115に従ってそのOSD60が処理しなければならないチャンクのリストを維持する。ノード10は、コピー・リスト115のチャンクを順番にコピーする。コピー・リスト115は、過度に大きいコピー・リストによってノードに負担をかけることを避けるために、おおよそ、CG毎に10Kのエントリを下回ることが好ましい。ステップ226において、ノード10が、各チャンクに関するフラグ・テーブル120を生成する。

【0060】

すべてのノード10は、それらのノード10に記憶される各チャンクに関してメモリ50に記憶されるデータ構造であるフラグ・テーブル120を管理する。クラスタ内の各ノードは、クラスタに記憶される各CGの各チャンクに関するフラグ・テーブルにアクセスすることができる。図12は、フラグ・テーブル120の例を示す。フラグ・テーブル120は、各チャンクが設定されたCLEANフラグ、DIRTYフラグ、またはDEGRADEDフラグを有するかどうかに関する情報を含む。加えて、フラグ・テーブル120は、DEGRADEDフラグが設定された後のイベント番号と、そのイベント番号の後にデータが書き込まれたイベント番号とを記憶する。設定されるフラグは、値「1」を有し、一方、値「0」はフラグが設定されなかったことを示す。ステップ227において、ノード10が、コピーされる各チャンクに関して、DIRTYフラグを設定する。ノード10がクラッシュする場合、リカバリを必要とするチャンクのリストが、DIRTYと印を付けられたチャンクに関してCGをスキャンすることによって再生成され、これは、下でより詳細に説明される。

【0061】

ステップ228において、ノード10が、そのノード10が受信したチャンクのリストの各チャンクに関するスタブを生成する。図13は、スタブで管理されるチャンク・スタブ管理情報105の例である。各OSD60は、BEING__CONSTRUCTEDと印を付けられた各CGの各チャンクに関してメモリ50にチャンク・スタブ管理情報105を記憶する。チャンク・スタブ管理情報105は、スタブが対応するチャンク番号を示すチャンク・フィールド106と、チャンクが対応するCGを示すCGIDフィールド107とを含む。チャンク・スタブ管理情報105は、チャンクの古いマップの位置108と、このチャンクに関するコピー・プロセスが始まったか否かを示すRECOVER_Y__READ要求送信済みフィールド109とをさらに含む。コピー・プロセスが始まった場合、RECOVER_Y__READ要求送信済みフィールドは、「1」で印を付けられ、そうでない場合、そのRECOVER_Y__READ要求送信済みフィールドは、「0」で印を付けられる。また、チャンク・スタブ管理情報105は、チャンクのすべてのスライスがコピーされたかどうかを全スライス・コピー済みフィールド110で示し、I/Oキー

10

20

30

40

50

・リストが空であるかどうかをキー・リスト空フィールド 1 1 1 で示す。O S D 6 0 は、すべてのスライスがコピーされたかどうかを確認するためにビットマップ 7 5 を調べる。すべてのスライスがコピーされた場合、値が「1」に設定され、そうでない場合、値は「0」である。O S D は、I / O キー・リスト 9 5 が空であるか否かを示すためにチャンクの I / O キー・リスト 9 5 にいずれかの I / O キーがあるかどうかを判定する。I / O キー・リスト 9 5 が空である場合、フィールドは「1」で印を付けられる。

【0062】

図 1 1 のステップ 2 2 9 において、ノード 1 0 が、各チャンクに関するスライスのビットマップ 7 5 を構築する。図 1 4 は、ビットマップ 7 5 の例である。ビットマップ 7 5 は、揮発性であり、メモリ 5 0 に記憶される。たとえば、3 2 K のスライスを用いると、6 4 M B のチャンクを記載するビットマップは、2 0 4 8 ビットまたは 2 5 6 バイトになる。ステップ 2 3 0 において、新しい位置が、古いマップ 7 0 のメンバーのノード 1 0 からソースの位置（ダウンしていないボリューム）を選択する。新しい位置は、C G I D 7 7 に関して C R U S H によって生成されたディスクのリスト 7 6 からソース・ボリューム（第 1 のダウンしていないボリューム）6 5 を選択する。たとえば、C R U S H は、古いマップ 7 0 からソース O S D 6 0 を選択する際にある程度の擬似ランダム性を提供し、すべての新しいマップ 7 1 のターゲットがそれらのターゲットの R E C O V E R Y _ R E A D 要求を同じソースに送信するように実装され得る。このプロセスは、基礎を成すファイルシステムのキャッシュを利用して、クラスタの処理の負荷に対するリカバリの全体的影響を小さくする。

10

20

【0063】

図 3 に示されるように、各 C G は、チャンクから構成され、さらに、各チャンクは、データのスライスから構成される。データのマイグレーション中、各チャンクは、チャンクの構成要素のスライスを転送することによってマイグレーションされる。ステップ 2 3 5 において、O S D 6 0 が、スライスを転送するために、ステップ 2 3 0 において決定されたダウンしていないボリュームに R E C O V E R Y _ R E A D 要求を送信する。R E C O V E R Y _ R E A D 要求は、ステップ 2 2 5 において受信され、メモリに記憶されたチャンクのリスト 1 1 5 にある各チャンクの隠すスライスに関して送信される。ステップ 2 3 6 において、チャンク・スタブ管理情報 1 0 5 の R E C O V E R Y _ R E A D 要求送信済みフィールド 1 0 9 が、そのチャンクに関するコピーが始まったことを示すために「1」に変更される。

30

【0064】

チャンクの各スライスが転送され、新しいノードに書き込まれるとき、ステップ 2 4 5 において、そのスライスに対応するビットが、そのスライスが書き込まれたことを示すためにビットマップ 7 5 内で変更される。ビットマップ 7 5 は、各チャンクのどのスライスが書き込まれたかを示す。チャンクに関する各スライスは、ステップ 2 4 0 において、ビットマップ 7 5 に従って順番に R E C O V E R Y _ R E A D 要求を用いてコピーされ、受信ノード 1 0 の新しい位置のディスク 6 5 に書き込まれ得る。図 1 4 に示される例示的なビットマップ 7 5 においては、マイグレーション中、スライス 1 が正常に転送された後、概して、残りのスライスに関する順番に従って、スライス 2 が次に転送される。図 1 4 のビットマップ 7 5 は、各スライスに対応する複数のビットを示す。少なくとも 1 つのビットが、スライスがコピーされ、ターゲット・ノードに書き込まれたか否かを表す。

40

【0065】

ノード 1 0 は、ビットマップ 7 5 を調べて、それぞれのチャンクのすべての内容がコピーされたかどうかを判定する（2 5 0）。ノードは、チャンク・コピー・リスト 1 1 5 の各チャンクに関するビットマップ 7 5 がすべてのスライスがコピーされたことを示すかどうかを判定する。ステップ 2 5 1 において、O S D が、そのチャンクに関してすべてのスライスがコピーされたことを示すためにチャンク・スタブ管理情報 1 0 5 の全スライス・コピー済みフィールド 1 1 0 を「1」に設定する。いずれかのスライスがコピーされていない場合、ステップ 2 3 5 において、それらのスライスに関して R E C O V E R Y _ R E

50

A D 要求が送信される。ステップ 2 5 2 において、チャンクが、ターゲット・ノードのチャンク・コピー・リスト 1 1 5 から削除される。ノードは、チャンク・コピー・リスト 1 1 5 の各チャンクに関するスライスのコピーを続ける。ステップ 2 5 3 において、すべてのチャンクがコピーされ終わっていない場合、プロセスが、ステップ 2 2 8 に戻る。すべてのチャンクの内容のすべてがコピーされたとき、新しい位置（すなわち、新しいノード）は、ステップ 2 5 5 において、フェーズ 2 に肯定応答する A C K メッセージを送信する。フェーズ 2 は、すべてのノードがそれらのノードのマイグレーション作業を完了したという報告を P a x o s に従って返すときに終了する。

【 0 0 6 6 】

ステップ 2 2 7 において、チャンクがたとえば W R I T E オペレーションによって修正されていること、またはチャンクがコピーされるべきであることを示すために、チャンクが、フラグ・テーブル 1 2 0 において D I R T Y フラグによって印を付けられる。D I R T Y フラグは、チャンクが書き込まれるかまたはそれ以外の方法で修正される前に設定される。D I R T Y フラグは、S A 5 5 が W R I T E オペレーションを送ったすべての O S D 6 0 が W R I T E に応答した後に、たとえば、S A 5 5 によって非同期にクリアされる。チャンクは、D R I T Y と印を付けられる場合、チャンク・コピー・リスト 1 1 5 上か、ビットマップ 7 5 上か、I / O キー・リスト 9 5 上か、またはビットマップ 7 5 と I / O キー・リスト 9 5 との両方の上かのいずれかにある。各チャンクのリカバリが完了されるとき、そのチャンクの D I R T Y フラグがクリアされ、このことは下で説明される。

【 0 0 6 7 】

マイグレーション中、S A 5 5 は、I / O オペレーションを受け取る場合、C G を記憶する古いマップ 7 0 の下のノード 1 0 と、C G を記憶する新しいマップ 7 1 の下のノード 1 0 との両方に、C G 内の特定のデータを対象とする I / O オペレーションを送らなければならない。このプロセスは、マイグレーションするチャンクの内容が、古いマップ 7 0 の下のノード 1 0 へのたとえばホスト 2 0 からの、S A 5 5 によって処理された W R I T E 要求によって上書きされる場合、チャンクの完全性が維持されることを保証する。

【 0 0 6 8 】

図 1 5 および 1 6 は、本発明の実施形態によるマイグレーション中の W R I T E 要求のプロセスの流れ図を示す。図 1 5 のステップ 4 0 0 において、ノード 1 0 が、W R I T E 要求を受信する。ステップ 4 0 5 において、受信ノード 1 0 の S A 5 5 が、現在の（古い）マップ 7 0 の下で記憶される C G に W R I T E が影響を与えるかどうかを調べる。S A 5 5 は、S A 5 5 が記憶する C G に W R I T E が影響を与えると判定する場合、ステップ 4 1 0 において、マップ・マイグレーション状態を調べる。S A 5 5 が記憶する C G に W R I T E が影響を与えないと S A 5 5 が判定する場合、プロセスは始めに戻る。ノード 1 0 は、マップ・マイグレーション状態にない場合、W R I T E 要求を O S D 6 0 に送信し、W R I T E が通常通り適用される（4 2 5）。ノード 1 0 がマップ・マイグレーション状態にあり、W R I T E 要求に関する C G が影響を与えられる場合、S A 5 5 は、古いマップ 7 0 と新しいマップ 7 1 との両方の下の C G を記憶するノード 1 0 に W R I T E 要求を送信しなければならない。ノードは、それぞれ、どの C G がマイグレーション状態にあるかを特定することができる。ステップ 4 1 5 において、ノードが、そのノードが古いマップ 7 0 と新しいマップ 7 1 との両方の下でそのノードのローカル・ストレージ・デバイス 6 7 に記憶された C G のコピーを有するかどうかを判定する。ノードが C G のコピーを記憶する場合、ただ 1 つの W R I T E 要求が O S D 6 0 に送信され、W R I T E が適用される（4 2 5）。W R I T E が通常通り適用され、処理が図 1 7 のステップ 8 0 0 に進む。そうではなく、ノードが古いマップ 7 0 と新しいマップ 7 1 との両方の下で記憶された C G のコピーを持たない場合、W R I T E 要求は、ステップ 4 2 0 において、新しいマップの下の C G を記憶するノードに送信され、処理は、下で説明される図 1 6 に示されるように進む。

【 0 0 6 9 】

図 2 6 は、I / O オペレーションの受け取りを扱うためのホスト 2 0 とノード 1 0 a お

10

20

30

40

50

よび 10b との間の通信フローを示す。図 26 において、ホスト 20 は、クラスタに記憶されたデータに対する I/O オペレーションの I/O 要求 2501 をクラスタ内のノード 10a に送信する。場合によっては、ホスト 20 から I/O 要求を受信するノード 10a が、I/O 要求に関するデータを記憶しない可能性がある。図 26 に示されるように、そのような場合、受信ノード 10a は、I/O オペレーションのための適切なノード 10b を指定するリダイレクト 2502 をホスト 20 に送信することによって適切なノード 10b に I/O 要求をリダイレクトし得る。リダイレクト 2502 に応答して、ホスト 20 は、I/O 要求を扱う適切なノードであるノード 10b に I/O 要求 2503 を送信し得る。I/O オペレーションを完了した後、ノード 10b は、応答 2504 をホスト 20 に送信し得る。

10

【0070】

図 27 は、クラスタ内のノード 10a へのクラスタに記憶されたデータに対する I/O オペレーションに関する代替的な通信フローを示す。図 27 において、ホスト 20 は、クラスタに記憶されたデータに対する I/O オペレーションの I/O 要求 2601 をクラスタ内のノード 10a に送信する。場合によっては、受信ノード 10a は、I/O オペレーションのための適切なノード 10b にリダイレクト 2602 を送信することによって適切なノード 10b に I/O 要求を直接リダイレクトし得る。I/O オペレーションを完了した後、ノード 10b は、応答 2603 をホスト 20 に送信し得る。

【0071】

図 17 は、通常の条件の間に I/O オペレーションを受け取るノード 10 のプロセスの流れ図を示す。通常の条件は、マップ・マイグレーション状態が 0 であり、マップ変更プロセスの下でのマイグレーションがノードで現在処理されていないことを示すときとして定義され得る。ステップ 800 において、ノード 10 に関する SA55 が、I/O オペレーションを受け取る。

20

【0072】

I/O オペレーションを受け取ると、受信ノードが、ノードに関する OSD60 に I/O オペレーションを送り (801)、OSD60 が、I/O オペレーションに対応するチャンクに関する I/O キー・リスト 95 にオペレーションを追加する (802)。ステップ 805 において、ノードが、チャンクに関するフラグ・テーブル 120 の DIRTY フラグを「1」に設定する。I/O オペレーションが、ステップ 810 において、チャンクに対して OSD60 によって実行される。ステップ 815 において、オペレーションが、対応するチャンクの I/O キー・リスト 95 から削除される。プロセスは、チャンクの I/O キー・リスト 95 にいずれかのその他のオペレーションが存在するかどうかをノード 10 が判定するときステップ 820 に続く。チャンクに関するその他のオペレーションがある場合、プロセスはステップ 810 に進み、オペレーションが実行される。いかなるその他のオペレーションもない場合、チャンクに対応するフラグ・テーブル 120 において DIRTY フラグが「0」に設定される (825)。代替的に、プロセスは、ステップ 825 に進む前にステップ 820 においていずれかのその他のオペレーションがあるかどうかを調べることなくステップ 815 からステップ 825 に進む可能性がある。

30

【0073】

新しいマップ 71 のメンバーのいずれかが WRITE 要求を完了することができない場合、SA55 は、新しいマップ 71 のチャンクに、フラグ・テーブル 120 の DEGRADED フラグによって DEGRADED として印を付ける。加えて、古いマップ 70 のノードの SA55 は、古いマップ 70 のいずれかが WRITE に失敗する場合、チャンクに DEGRADED として印を付ける。チャンクは、チャンクのいずれかのコピー (いずれかの CG のコピー) がオンラインおよび応答可能ではないことがノード 10 によって知られているとき、チャンクが書き込まれていることを示すために DEGRADED として印を付けられる。チャンクは、SA55 が I/O 要求に DEGRADED 状態を生じているものとしてフラグを立てるためか、または SA55 が丁度書き込まれた DIRTY チャンクの DEGRADED の印を要求し、いくつかのミラー (たとえば、CG のコピーを記憶

40

50

するノード)がWRITE中にエラーを返したかもしくはクラッシュしたが、それでも、高い信頼性でデータをリカバリするのに必要な最小限の数を形成する十分なCGのコピーが存在することを知っているためかのどちらかでDEGRADEDとして印を付けられる。チャンクにDEGRADEDとして印を付けることは、DIRTYフラグを削除することと、アクセス不可能なコピーがアクセス可能になるとき、たとえば、ノードが障害の後にオンラインに復帰するときにチャンクをリカバリする必要性に関する情報を失わないことを可能にする。加えて、古いマップに対応するノードのSA55は、古いマップ70のいずれかがWRITEに失敗する場合にチャンクにDEGRADEDとして印を付ける。古いマップに対応するノードのすべてのノード10は、そのノード10がCGのすべてのコピーが存在し、リカバリされると判定する場合、DEGRADEDフラグをクリアする。

10

【0074】

マイグレーション・プロセス全体を通じて、新しいマップ71の下でデータを記憶しないノード10を含む、古いマップ70のすべてのノード10は、I/O要求を受け付け続ける。古いマップ70のノードの下での各ノード10に関するSA55は、図15を参照して上で説明されたように、WRITE、TRUNCATE、およびREMOVE要求を新しい位置へ転送する。SA55は、古い位置への要求と平行してそれぞれの新しい位置にオペレーションを送る。

【0075】

図16は、古いマップの下でのCGのコピーを記憶するSA55からのWRITE要求を受信するターゲット・ノードのプロセスの流れ図を示す。ステップ435において、新しい位置が、WRITEを受信する。ステップ440において、ノード10が、全スライス・コピー済みフィールド110を参照することによってチャンクのすべてのスライスがコピーされたかどうかを判定するためにチャンク・スタブ管理情報105を調べる。チャンクは、ソースから既にコピー済みである(つまり、そのチャンクがCLEANである)場合、ステップ445においてDIRTYとして印を付けられ、WRITEが、チャンクに直ちに適用される(ステップ450)。チャンクがDIRTYとして印を付けられる場合、OSD60は、ステップ455において、ビットマップ75を調べて、WRITEが当てはまるスライスがコピー済みであるか否かを判定する。ステップ455において、ビットマップ75がスライスがコピー済みであることを示す場合、ステップ490において、WRITEが、対応するスライスに適用される。書き込みが当てはまるスライスがコピー済みでない場合、ステップ460において、ノードが、WRITEオペレーションによって書き込まれるべきであるデータがスライス・サイズ未満であるかどうかを調べる。WRITEによって書き込まれるべきであるデータがスライス・サイズ以上である場合、WRITEは、適切なスライスの位置に適用される(480)。ビットマップ75のビットが、スライスがマイグレーション中に新しいマップの位置にコピーされることに晒されないようにそのスライスに関して変更される(ステップ485)。WRITEによって書き込まれるべきであるデータがスライス・サイズ未満である場合、そのスライスに関するデータがマイグレーションされ得るように、そのスライスに関するRECOVERY__READ処理が、ビットマップ75の順序を外れて(つまり、スライスの順序を外れて)送られる(470)。つまり、そのスライスに関するRECOVERY__READが、ビットマップ75による順番で通常は次に転送されるスライスの前に送られる。そのスライスに関するデータがマイグレーションされ、コピーされる(475)と、WRITEオペレーションが適用され、ステップ485において、ビットマップ75のビットが、そのデータがコピーされたことを示すためにそのスライスに関して変更される。

20

30

40

【0076】

(トランケートするための)TRUNCATEまたは(削除するための)REMOVEなどのその他の修正オペレーションが、古いマップ70と新しいマップ71との両方の下のノード10に送られる。概して、WRITE要求に関する同じ処理が、TRUNCATEおよびREMOVEオペレーションに当てはまる。トランケート処理は、コピー・プロ

50

セスがどこまで進捗しているかに応じて異なる。チャンクがまだコピーされていない場合、TRUNCATEオペレーションは、成功する非オペレーション(non-operation)である。チャンクは、既に完全にコピー済みである(そのチャンクがCLEANである)場合、フラグ・テーブル120においてDIRTYと印を付けられ、TRUNCATEオペレーションが、通常通り適用される。TRUNCATEオペレーションがあるときにチャンクがDIRTYであり(コピーされており)、ビットマップ75が、トランケート・ポイント(truncation point)が現在のコピー・オフセット(copy offset)を過ぎている(まだコピーされていない)ことを示す場合、TRUNCATEオペレーションは、成功する非オペレーションである。TRUNCATEオペレーションがあるときにチャンクがDIRTYであり、ビットが、チャンクのこのスライスが既にコピー済みであることを示す場合、チャンクは、ソースからのさらなるコピーを防ぐためにロックされる。そして、チャンクは、TRUNCATEオペレーションによってトランケートされ、ソースからさらにコピーされない。

【0077】

マイグレーション中に、I/OオペレーションがREMOVEである場合、チャンクが、新しい位置から削除され、インメモリ構造からパージされる。この場合、OSD60は、チャンクが削除される前に、もしあれば、チャンクの保留中のコピーがコピーされるまで待つ。チャンクがCGから削除されている場合、新しいマップのメンバーは、そのチャンクも削除する。

【0078】

マイグレーション中のI/O処理中に、新しい位置のOSDは、通常のWRITEオペレーションにおいてOSD60が行うのと同じ方法で要求元のSAのキーが追跡され続けるようにする。図18は、本発明による、マイグレーション中のI/Oオペレーションが完了した後にCHUNKをクリーンに設定するプロセスの流れ図を示す。ステップ505において、オペレーションが完了されたとき、OSD60が、SAにACKメッセージによって応答する(ステップ510)。その後、OSD60は、古いマップ70と新しいマップ71との両方のOSDにおいてノード10にCLEAN指令を送信する(520)。ステップ525において、各チャンクに関して、それぞれの新しい位置が、チャンクのI/Oキー・リスト95からI/Oキーを削除する。新しい位置のノード10は、チャンクをCLEANに設定する前に次のこと、すなわち、チャンクの残りのI/Oキーのリスト95が空であるかどうか(ステップ530)、ビットマップ75のビットがすべてのスライスがコピー済みであることを示すかどうか(ステップ535)、およびチャンクがチャンク・コピー・リスト115にあるかどうか(540)を調べる。代替的に、ノード10は、ステップ535において、チャンク・スタブ管理情報105を参照して全スライス・コピー済みフィールドが「1」であるかまたは「0」であるかを判定することによって、すべてのスライスがコピー済みであるかどうかを調べる可能性がある。これらの要因のいずれかが満たされない場合、チャンクは、ステップ550においてフラグ・テーブル120でDIRTYのままにされる。これらの要因のすべてが満たされる場合、DIRTYフラグが、そのチャンクに関してクリアされ、チャンクは、チャンクに対応するフラグ・テーブル120でステップ545においてCLEANとして印を付けられる。ステップ555において、ノードが、チャンクがクリーンであると印を付けられるために必要なプロセスを実行する。たとえば、ステップ530において、チャンクのI/Oキー・リストが空でないと判定される場合、ノードは、残りのI/Oオペレーションを実行する。プロセスが完了した後、チャンクが、チャンクをCLEANに設定するために再び調べられる(つまり、ステップ530~540)。チャンクをCLEANに設定すべきかどうかを判定する際に考慮される可能性がある要因は、少なくとも最小限の数を示す予め決められた値に関するCGに対するWRITEオペレーションが正常に完了されたかどうかを判定することを含む。チャンクは、DIRTYのままにされ、最小限の数があるときにリカバリされ得る。この場合、すべてのCGのコピーおよびそれらのCGのコピーのチャンクが、Paxosのイベント番号を考慮に入れることを含めて調べられる可能性がある。換言すると

10

20

30

40

50

、システムは、十分なコピーが利用可能であるとき、チャンクに関するすべてのコピーが同一であることを確認する。代替的に、D I R T Yフラグは、少なくともC Gのコピーの最小限の数に対するW R I T Eオペレーションが正常に完了されたときにクリアされる。最小限の数は、チャンク／C GのC L E A Nなコピーを有するノードの最少数である。

【 0 0 7 9 】

最小限の数が確立されると、すべてのノードが、チャンクのいずれかのコピーがD I R T Yフラグを有するかどうかを判定する。D I R T Yフラグを有するチャンクがある場合、ノードに記憶されたD I R T Yなコピーを有するノードが、ダウンしていないボリュームを選択し、チャンクをリカバリするためにR E C O V E R Y _ R E A D要求を送信する。ノードは、新しいマップ7 1からダウンしていないボリュームのメンバーを選択する。

10

【 0 0 8 0 】

ノードのクラッシュ

マイグレーション中、ノード1 0またはノード1 0のディスク6 5がクラッシュする可能性があり、それがマイグレーションを妨げ、この場合、マイグレーション中に継続しているW R I T Eオペレーションは、ノードがクラッシュする前に処理されるチャンクの一部に過ぎない可能性があり得る。この状況は、そのようなチャンクがフラグ・テーブル1 2 0でD I R T Yであり、C Gがマイグレーション中にB E I N G _ C O N S T R U C T E Dとして印を付けられ、このことが、対応するデータの内容が信頼できないことを示すために検出され、リカバリされ得る。

【 0 0 8 1 】

図1 9は、ノードがクラッシュの後にオンラインに復帰するときに従うプロセスに関する流れ図を示す。ノード1 0またはディスク6 5がオンラインに復帰するときに、ノードは、ステップ5 6 0において、どのチャンクがD I R T Yとして印を付けられているか、およびどのC GがB E I N G _ C O N S T R U C T E Dとして印をつけられているかを判定する。代替的に、ノード1 0は、ノードのクラッシュまたは障害の後にリストが利用可能である場合、チャンク・コピー・リスト1 1 5を参照して、どのチャンクがコピーされる必要があるかを判定し得る。この判定から、ノード1 0は、ステップ5 6 5において、コピーされるチャンクのリストを生成する。ノードは、一部のスライスがクラッシュの前にコピーされたものとして示されているとしても、コピーされるチャンクに関してビットマップ7 5のビットをクリアする(5 7 0)。加えて、ステップ5 7 0において、チャンク・スタブ管理情報1 0 5がクリアされる。ノード1 0は、ステップ5 7 5においてチャンクをコピーするために妥当なソース・ノード(ダウンしていないボリューム)を選択し、新しく書き込まれるビットを含むチャンク全体が、R E C O V E R Y _ R E A D要求を用いて妥当なソースからコピーされる。ステップ5 8 0において、処理が、図1 1のステップ2 3 5に進んで、ステップ5 7 5からの決定されたダウンしていないボリュームにR E C O V E R Y _ R E A D要求を送信し、データをコピーするためにフェーズ2のステップ2 3 5に続く処理に従って進行する。

20

30

【 0 0 8 2 】

ノードは、通常の実行中に(つまり、マップ変更またはマイグレーション中でないときに) クラッシュする場合、図1 1のステップ2 5 5のようにP a x o sにA C Kメッセージをポストしないことを除いて上記のルーチンに従って回復する。ノードは、W R I T Eオペレーションなどのデータを修正するI / Oオペレーション中にクラッシュする場合、決定されたダウンしていないボリュームから新しく書き込まれたデータをリカバリする。したがって、データの完全性が維持される。

40

【 0 0 8 3 】

妥当なソースは、マイグレーションを再開する古いマップ7 0からのノード、または妨げられたマップ変更を完了するために十分に自己をリカバリした新しいマップ7 1のメンバーからのノードである可能性がある。そして、それぞれの障害が起こったノードが、上述のように、チャンクをC L E A Nとして設定する処理を経る。

【 0 0 8 4 】

50

マップ・マイグレーション中のソース・ノードの障害

ソース・ノード 10 または ノード 10 のディスク 65 がマイグレーション中にクラッシュする場合、図 20 に示されるように以下の処理フローが開始される。ステップ 601 において、クラッシュの前に完全にコピーされなかったそれぞれの CG に関して、障害が起こったソース・ノードからチャンクを受信していた新しいマップ 71 のメンバーは、CRUSH を用いて、マイグレーションされる必要がある CG に関してそれに記憶された CG を有する古いマップ 70 の下でディスク・リスト 76 を生成する。そして、ノードは、ステップ 605 において、リストの第 1 のダウンしていないボリュームを選択する。第 1 のダウンしていないボリュームであるように決定されるボリューム 65 は、「有効な」コピーと呼ばれる。有効なコピーを保持すると決定されたノードは、（クラッシュし、今ではオンラインに復帰している同じノードを含む）古いマップ 70 と新しいマップ 71 との両方の下のピア・ノードをリカバリする。有効なコピーは、ここで、ターゲット・ノードから RECOVERY_READ を受信する（たとえば、ステップ 235）。言い換えると、ターゲット・ノードが、データ転送のソースとなるように古いマップのメンバーの中で新しいダウンしていないボリュームを選択する。別の言い方をすれば、上述の通常のリカバリ・プロセスは、同じである。有効な OSD は、マイグレーション中にクラッシュした 1 つのノードまたは複数のノード 10 がオンラインに復帰するときにそれらのノード 10 をリカバリする。有効な OSD は、すべてのノードが利用可能な状態情報 100 を調べることによって、クラッシュした位置がオンラインに復帰していることを知っている。

10

【0085】

20

本発明のこの態様によれば、ソース・ノード 10 が障害から復帰するために単一のルーチンのみが必要とされる。ルーチンは、セルのマップ・マイグレーション状態を調べて、ピア・リストを生成するとき、およびどの OSD が有効な OSD であるかを決定するときに、それらのセルがマップ・マイグレーションのフェーズ 2 にあるかどうかを判定する。

【0086】

書き込み不可能と判定される CG

ターゲット・ノード 10 がクラッシュするか、またはターゲット・ノード 10 のディスク 65 がクラッシュする場合、CG のコピーを記憶するノード 10 が、ターゲット・ノードの障害があっても最小限の数が維持されるかどうかを判定する。最小限の数は、CG のコピーを記憶するノード 10 またはディスク 65 の数が最小限の数の値を超えているときに維持される。最小限の数の値は、各ノード 10 が知っている予め決められた値である。代替として、最小限の数は、CG に関するチャンクを含む OSD 60 の数が予め決められた最小限の数の値を超える場合にやはり満たされる可能性がある。新しいマップ 71 の下のノード 10 が最小限の数を維持する場合、リカバリが継続する。新しいマップ 71 のメンバーの数が最小限の数を下回る場合、残りの OSD 60 が、リカバリを完了するが、CG は、書き込み不可能と判定される可能性がある。CG が書き込み不可能であると判定される場合、この状態は、次のマップ変更があると変更される。CG は、最小限の数が現在利用可能でない場合、書き込み不可能であると判定される。

30

【0087】

ソース・ノード 10 がクラッシュするか、またはソース・ノード 10 のディスク 65 が障害にあう場合、マイグレーションは、（上で説明されたように）ダウンしていないその他のソース OSD 60 がある限り継続する。ソースのコピーのすべてがダウンしている場合、ターゲット・ノード 10（受信ノード）はマイグレーションを断念する。この場合、ターゲット上で構築された CG のコピーは、BEING_CONSTRUCTED と印を付けられたままであり、ソースのコピーがオンラインに復帰し、利用可能にされるときに可能であればリカバリされる。この場合、BEING_CONSTRUCTED として印を付けられた CG のコピーは、リカバリの目的を除いて読み取り可能または書き込み可能でない。

40

【0088】

フェーズ 3

50

図 2 1 は、本発明の実施形態によるフェーズ 3 のプロセスの流れ図を示す。ステップ 3 0 1 において、提案元が、すべてのノード 1 0 がステップ 2 5 5 において示されるように肯定応答を生成することによってフェーズ 2 に肯定応答したかどうかを調べる。すべてのノード 1 0 がフェーズ 2 に肯定応答したとき、フェーズ 2 は完了する。提案者は、提案されたマップ 7 1 のコミットである、マップ変更のフェーズ 3 を開始する。ステップ 3 0 5 において、マップ変更に関する M A P _ C O M M I T イベント・メッセージが、クラスタのすべてのノードに P a x o s に従って送信される。ステップ 3 1 0 において、コミット・メッセージを受信したすべての O S D が、それらの O S D が記憶する C G を評価する。ステップ 3 1 5 において、各 O S D が、その O S D が記憶する各 C G を新しいマップ 7 1 と付き合わせて、その O S D が記憶するあらゆる C G に関して、ノード 1 0 が新しいマップの下の C G の適切な保有者であるかどうかを調べる。

10

【 0 0 8 9 】

ノード 1 0 は、そのノード 1 0 が記憶する C G が新しいマップ 7 1 の下でノードによって記憶されるものとして示されていないと判定する場合、ステップ 3 2 0 において、永続的な C G _ I N V A L I D _ P L A C E M E N T フラグによって C G に印を付け、ステップ 3 3 0 において、その C G のチャンクのすべての非同期の削除および空間の再利用を開始する。その C G に関して記憶されたすべてのチャンクが、ノード 1 0 から削除される。チャンクの削除の後、ステップ 3 3 5 において、C G の階層およびチャンク・スタブ管理情報 1 0 5 が削除され、さらに、C G が O S D 6 0 に対応するキャッシュから削除される。O S D は、その O S D が新しいマップ 7 1 に従って C G を適切に記憶することを発見する場合、B E I N G _ C O N S T R U C T E D フラグが設定されている場合、ステップ 3 2 5 において、C G からそのフラグをクリアする。ステップ 3 4 0 において、新しいマップが、コミットされ、あらゆる S A が、今や新しいマップを使用し、あらゆるノードが、今や新しいマップ 7 1 のデータ配置に従って I / O オペレーションを受け付ける。

20

【 0 0 9 0 】

冗長性のポリシーの変更

当業者は、ストレージ・システムにおいて冗長性のレベルが変更され得ることを理解するであろう。たとえば、ストレージ・システムの冗長性のポリシーが、R A I D 1 ~ R A I D 6 の間で設定される可能性があり、その後、変更される可能性がある。冗長性のレベルは、下のようにクラスタに反映される可能性があり、R A I D 1 に関しては異なるミラー・カウント (m i r r o r c o u n t) を用いてポリシー変更に関して説明される。

30

【 0 0 9 1 】

ミラー (コピー) ・ポリシーは、たとえば、管理コマンドによって変更され得る。また、マップ変更の提案が、冗長性のレベルの変更に基づいて行われる可能性がある。ポリシーの変更が増加であるかまたは減少であるかに関わらず、新しいマップ 7 1 が生成される。ミラーの数を増やす場合、冗長な C G のコピーの数が新しいマップのレイアウトに従ってクラスタ内で配信されなければならないので、マップ変更が行われる。ミラーの数を減らす場合、クラスタ内のノードから C G のコピーを削除し、クラスタ内の残りのデータの配置を再配信するためにマップ変更が行われる。

40

【 0 0 9 2 】

図 2 2 は、各 C G に関連するミラー・カウントを示すポリシー・テーブル 8 5 の例である。この例において、ミラー・カウントは、各 C G に関して指定されるが、クラスタ全体またはクラスタのサブセクションに関して指定される可能性もある。ミラー・カウントは、たとえば、整数 3 であり、管理者によって指定される可能性がある。加えて、ミラー・カウントの既定値が、クラスタ内のすべての C G に関してストレージ・システムによって設定される可能性がある。

【 0 0 9 3 】

C R U S H のシーケンスに従って、ディスク 6 5 (ボリューム) が、シーケンスの終わりに追加および削除される。つまり、C R U S H は、3 ディスクに関して求められる場合、ディスク・リスト 7 6 { A , B , C } を生成する。そして、C R U S H は、4 ディスク

50

に関して求められる場合、ディスク・リスト 76 { A , B , C , D } を生成し、一方、2 ディスクに関して求めることは、ディスク・リスト 76 { A , B } を生成する。変更の性質に関係なく、ストレージ・システムは、リストの第 1 のダウンしていないメンバーが前のディスク・リスト 76 にあったか、またはそうではなく、データが完全にアクセス不可能であり、前の CG のリストからのディスク 65 が利用可能になるまでチャンクに関して何も行われ得ないかのどちらかを保証される。

【 0 0 9 4 】

図 23 の流れ図は、ストレージ・システムにおけるポリシー変更のためのプロセスを示す。管理者は、CG に関するミラー・カウントを増やすかまたは減らす。新しいミラー・カウントは、ポリシー・テーブル 85 に記憶される。ポリシー・変更が行われるために、POLICY__CHANGE イベント・メッセージが、Paxos に従ってクラスタ内のノード 10 に送信される。POLICY__CHANGE イベント・メッセージは、ポリシー変更が CG に関するミラー・カウントの増加であるかまたは減少であるかを含む。メッセージは、ステップ 701 において、クラスタ内の OSD 60 によって受信される。ステップ 705 において、ノードが、そのノードが記憶するポリシー変更によってどの CG が影響を与えられるかを判定する。ステップ 710 において、影響を与えられる各 CG について、ノードが、変更に基づいて新しいマップ 71 を構築する。ポリシー変更が増加である場合、それぞれの CG のレプリカ・コピーに関する新しい位置を含む新しいマップ 71 が構築される。たとえば、ポリシー変更がミラーの数を 2 増やした場合、新しいマップ 71 は、CG のレプリカ・コピーに関して 2 つの新しい位置を有する。新しいマップに従って、ポリシー変更に合致するように生成された CG の新しいコピーが、ポリシー変更の前に CG を記憶しなかったクラスタ内の位置に記憶される。ポリシー変更が減少である場合、新しいマップ 71 は、古いマップ 70 の下で含まれた CG の (古いポリシーと新しいポリシーとの間のミラー・カウントの差に応じた) 1 つまたは複数のレプリカ・コピーを含まない。

【 0 0 9 5 】

ステップ 715 において、ノードが、ポリシー変更が CG のコピーの数を減らすかまたは CG のコピーの数を増やすかを判定する。ポリシーが、コピーの数を減らし、OSD 60 が、ステップ 720 において、その OSD 60 が新しいマップの下で CG のコピーを記憶するようにもはや求められないことを発見する場合、その OSD 60 は、ステップ 725 において、その OSD 60 のローカル・ストレージからの CG の非同期の削除を開始する。ステップ 720 中に、ノードが、そのノードが新しいマップ 71 の下の CG のコピーを記憶すると判定する場合、たとえその他のコピーがクラスタのその他の場所で削除される可能性があるとしても、そのノードは、図 9 のステップ 135 に進む (つまり、フェーズ 1 で Paxos に受け付けをポストする) 。

【 0 0 9 6 】

しかし、ポリシーがコピーの数を増やす場合、現在の (古い) マップ 70 からのノード 10 の OSD 60 が、ディスク・リスト 76 が古いポリシーの下で利用可能であるかどうかに応じて、新しいマップ 71 の下の CG を記憶するすべてのメンバーか、または新しいマップ 71 の下の新しいノード 10 のみかのどちらかに、それらのノードが今や CG のコピーを記憶しなければならないことを知らせる。ステップ 740 において、CRUSH によって生成されたディスク・リストの第 1 のダウンしていないボリュームが、新しいマップ 71 の下のすべてのターゲット・ノードのメンバーにそれらのメンバーが記憶すべき CG を有することを知らせるノード 10 であるように決定される。その後、新しいメンバーが、新しいマップ 71 の下の各ターゲット・ノード 10 に CG__MIGRATION__REQUEST を送信することによってステップ 745 において通知される。ステップ 750 において、プロセスは、マップ変更プロセスのフェーズ 1 のステップ 150 に進む。

【 0 0 9 7 】

リストがポリシー変更のために生成されるときに利用可能でないディスク 65 に特定の CG が記憶される場合、そのような CG は、ダウンしていないいかなるノードによっても

見つけられず、ポリシー変更のリカバリを経ない。しかし、ポリシー変更イベントは、記録され、ターゲット・ノード 10 がオンラインに復帰し、ノードがその時点でリカバリ・プロセスを経るときに C G を記憶する O S D 6 0 に P a x o s に従って利用され得る。

【 0 0 9 8 】

本発明のストレージ・システムは、ディスク 6 5 およびノード 1 0 の障害を追跡する。予め決められた期間の後、オフラインになった（１つもしくは複数の）ディスク 6 5 または（１つもしくは複数の）ノード 1 0 がオンラインに復帰しない場合、それらのディスク 6 5 またはノード 1 0 は、マップ構成から削除され、新しいマップ 7 1 が、提案される。この場合、新しいマップ 7 1 が、上述のマップ変更プロセスによって、データ・マイグレーション、および新しいマップ 7 1 へのデータの複製をトリガする。結果として、冗長性のポリシーによって指定された冗長性のレベルがリカバリされる。予め決められた時間は、何らかのデータのすべてのコピーの全損失をもたらすその後の障害の可能性を最小化することの考慮に基づく可能性がある。また、システム管理者が、マップ変更コマンドを発する可能性がある。

【 0 0 9 9 】

部分的に書き込まれたデータのリカバリ

以下のルーチンに関するステップの一部の詳細は、既に説明されており、以下で繰り返されない。各ノードが、クラスタ内のイベントの通知を受信する。イベントは、ディスクもしくはノードの障害、またはディスクもしくはノードが再び加わることもしくは追加されることである可能性がある。発生が通知されるとクラスタの各ノードが、そのノードが記憶する C G を評価して、そのノードがイベントによって影響を与えられる C G のいずれかのコピーを記憶するかどうかを判定する。そのとき、ノードは、発生に基づいて影響を与えられる C G に関するチャンクを再構築またはコピーするためのアクションをそのノードが行う必要があるかどうかを判定する。図 2 4 は、本発明の実施形態による、クラスタに再び加わった後の部分的に書き込まれたデータのノードのリカバリのプロセスの流れ図を示す。図 2 5 は、本発明の実施形態による、クラスタに再び加わった後の部分的に書き込まれたデータのノードのリカバリのプロセスのメッセージ・シーケンスを示す。

【 0 1 0 0 】

ステップ 9 0 0 において、ノードまたはディスク 6 5（論理ボリューム）が、そのノードまたはディスク 6 5 がオフラインであった後にクラスタに再び加わる。ディスクが I / O オペレーションに失敗し、再試行またはリカバリ I / O オペレーションも失敗した場合、ディスクはオフラインであると宣言され得る。ノード / ディスク（ボリューム）は、オンラインに復帰し、アクティブ状態であり、ノード 1 0 は、再び加わったノードと呼ばれる。すべてのノードは、通知、および状態情報 1 0 0 の参照によってノード 1 0 がクラスタに再び参加したことを知っている（9 0 1）。ステップ 9 0 5 において、再び加わったノードを含むすべてのノードが、現在のマップの下で、そのノードが再び加わったノードに記憶されるいずれかの C G を有するかどうかを判定する。そのノードが判定する各 C G が再び加わったノードにコピーを有するので、ノードは、フラグ・テーブル 1 2 0 を調べることによって、C G に関するチャンクのいずれかが設定された D E G R A D E D または D I R T Y フラグを有するかどうかを判定する。

【 0 1 0 1 】

ステップ 9 1 0 においてはと判定すると、ステップ 9 1 5 において、ノード 1 0 が、設定された D E G R A D E D または D I R T Y フラグを有するチャンクのリスト 1 1 5 を計算し、再び加わったノードにリストを送信する。再び加わったノードは、ステップ 9 1 5 において、複数のノードからリストを受信する。ステップ 9 1 5 において各ノードによって送信されたリストは、各チャンクに関する情報を含む。情報は、フラグ・テーブル 1 2 0 から決定される。リストの情報は、たとえば、チャンクが設定された D E G R A D E D フラグを有するかまたは設定された D I R T Y フラグを有するか、そのイベント番号の後にチャンクへの W R I T E オペレーションによってチャンクが書き込まれたイベント番号、およびそのイベント番号の後に D E G R A D E D フラグが設定されたイベント番号を

含む。イベントは、Paxosからのメッセージであり、イベント番号は、クラスタにデバイスの障害を知らせるメッセージに関連するシーケンス番号である。したがって、イベントがディスクの障害である場合、メッセージは、WRITEが実行された番号よりも大きなシーケンス番号を用いてPaxosによって送信される。したがって、DEGRADED状態に関連するPaxosのシーケンス番号は、システムがどのコピーが古いかを曖昧性なく判定することを可能にする。

【0102】

915におけるリストの受信の後、再び参加したノードが、リカバリされる必要があるチャンクを有する各CGに関する現在のマップの下での第1のダウンしていないボリュームを決定する。ノードは、リストの情報に基づいて第1のダウンしていないボリュームを決定する。リストの情報から、ノードは、どのノードがチャンクの最も新しいコピーを記憶するか、およびどのコピーがCGに関するチャンクを再構築するのに十分であるかを判定し得る。ダウンしていないボリュームを決定すると、再び加わったノードは、リカバリされるかまたは再び書き込まれる必要がある各ノードに関してダウンしていないボリュームにRECOVER_Y__READ要求を送信する(925)。ステップ930において、その他のノード上のチャンクのその他のコピーに基づいて、再び加わったノードが、データを受信し、チャンクを再び書き込むか、またはチャンクを再構築する。ステップ935において、再び加わったノードが、現在のマップの下でそのノードが記憶するCGに関してデータのすべてがリカバリされたかどうかを判定する。再び加わったノードは、データのすべてをリカバリしていない場合、すべてのチャンクが再構築されるまでRECOVER_Y__READ要求を送信する。

10

20

【0103】

再び加わったノードは、ステップ935において、すべてのチャンクが再び書き込まれるかまたは再構築されると判定するとき、ステップ940において、情報のリストを送信した(914)すべてのノードにリカバリが完了していることを知らせるためにそれらのノードにメッセージを送信する。ステップ945において、すべてのノードが、チャンクのすべてのコピーが存在し、リカバリされるかどうかを判定する。すべてのチャンクは、チャンクのコピーを記憶する各ノードおよびディスクがオンラインである場合、ならびにすべてのチャンクのコピーが同じであり、そのDEGRADED状態に関連する最も大きなイベント番号を有するチャンクと同一であると判定される場合に存在し、リカバリされる。すべてのチャンクが存在し、リカバリされる場合、再び加わったノードを含むすべてのノードが、フラグ・テーブル120の各チャンクに関するDEGRADEDフラグをクリアする(ステップ955)。ステップ945において、ノードは、そのノードの判定中にそのノードのチャンクのいずれも存在せず、リカバリされないことを発見する場合、リカバリされる必要があるチャンクに関してダウンしていないボリュームを計算し、データをリカバリするためにダウンしていないボリュームにRECOVER_Y__READ要求を送信する。

30

【0104】

代替的な方法においては、すべてのノードが、ステップ905において、それらのノードが再び加わったノードにCGのコピーを記憶すると判定した後、各ノードは、影響を与えられるCGのリストおよび関連するメタデータを再び加わったノードに送信する。そして、再び加わったノードは、CGのリストを計算し、CGの独自のリストおよびメタデータを位置の受信されたリストのすべてのノードに送信する。CGのリストおよびメタデータを受信した後、再び加わったノードを含む各ノードは、受信されたCGのメタデータを調べ、そのノードが記憶するCGのコピーがその他のコピーに存在する欠けているデータであるかどうか、またはノードのコピーがその他のコピーからのリカバリを必要としているかどうかを判定する。ノードのコピーがその他のコピーからのリカバリを必要としている場合、ノードは、図24で説明されたルーチンのステップ920以降に従ってデータをリカバリする。

40

【0105】

50

本明細書において提供された以上の開示および教示から理解され得るように、本発明の実施形態は、ソフトウェアまたはハードウェアの制御論理またはこれらの組合せの形態で実装され得る。制御論理は、本発明の実施形態で開示された１組のステップを実行するように情報処理デバイスに指示するように適合された複数の命令として情報ストレージ媒体に記憶される可能性がある。本明細書において提供された開示および教示に基づいて、当業者は、本発明を実施するためのその他のやり方および／または方法を理解するであろう。

【 0 1 0 6 】

上記の説明は、例示的であり、限定的でない。本発明の多くの変形形態が、本開示を検討すると当業者に明らかになるであろう。したがって、本発明の範囲は、上記の説明を参照して決定されるべきでなく、その代わりに、添付の請求項をそれらの請求項の完全な範囲およびそれらの均等物と一緒に参照して決定されるべきである。

10

【 図 1 】

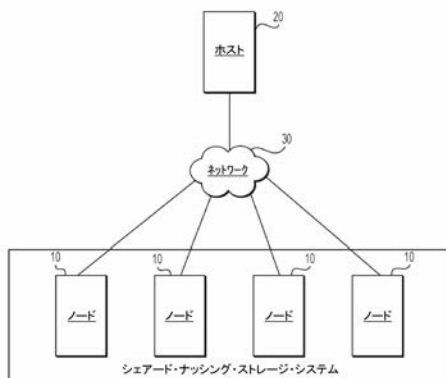


FIG. 1

【 図 2 】

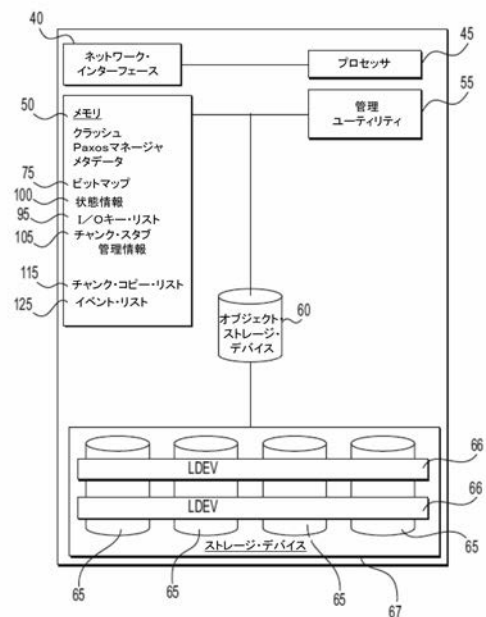


FIG. 2

【 図 3 】

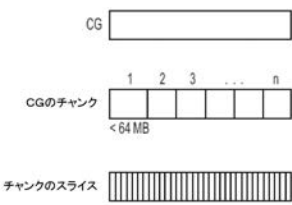


FIG. 3

【 図 4 】

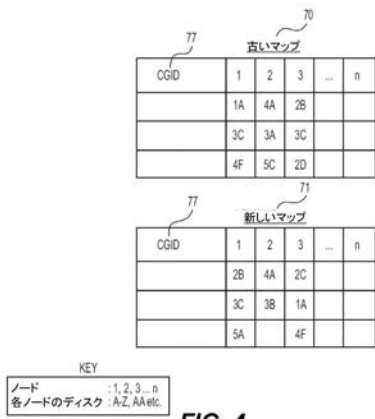


FIG. 4

【 図 5 】

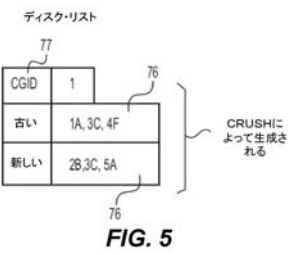


FIG. 5

【 図 6 】



FIG. 6

【 図 7 】

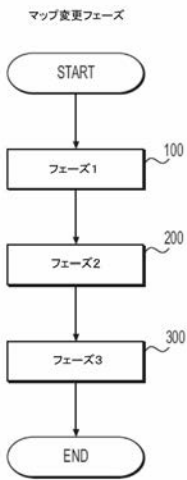


FIG. 7

【 図 8 】

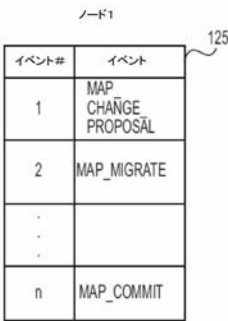


FIG. 8

【 図 9 】

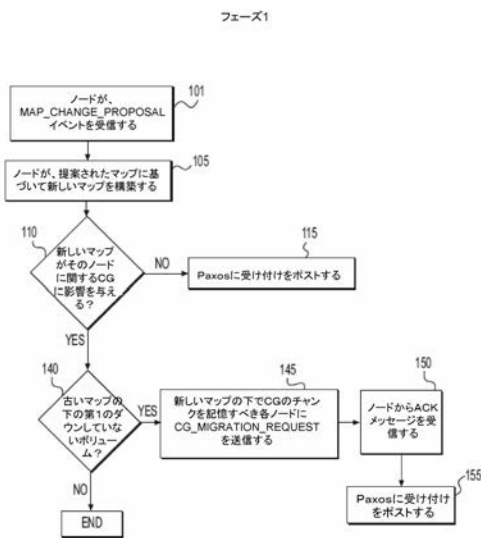


FIG. 9

【 図 1 0 】

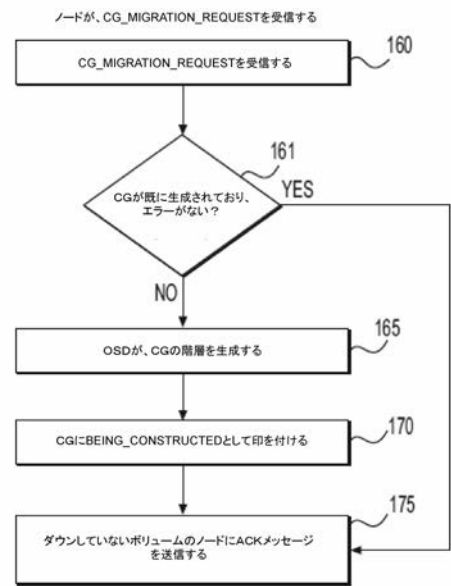


FIG. 10

【 図 1 1 】

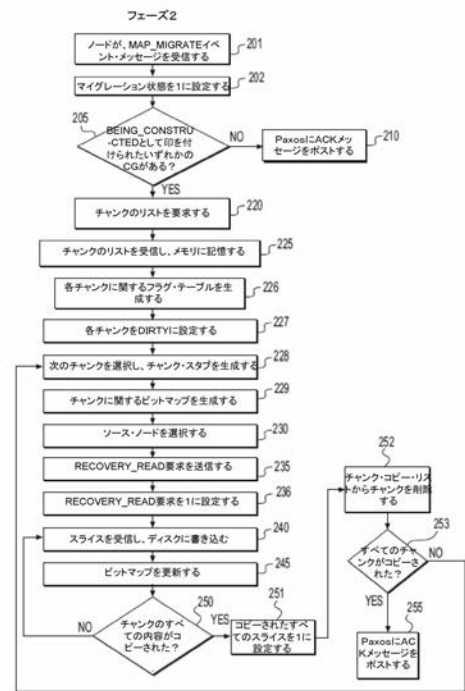


FIG. 11

【 図 1 2 】

フラグ・テーブル

	CHUNK 1		CHUNK 2	
	STATUS	DEGRADED フラグまたはW RITEの後のイ ベント番号	STATUS	DEGRADED フラグまたはW RITEの後のイ ベント番号
CLEAN	0		1	
DIRTY	1		0	
DEGRADED	1	3	1	2
WRITE		2		1

FIG. 12

【 図 1 3 】

チャンク・スタブ管理情報

106	チャンク番号	8	105
107	CG ID	4	
108	古いマップの位置	2A	
109	RECOVERY_READ 要求送信済み	0	
110	全スライス・コピー済み	0	
111	L/Oキーリスト空	0	

FIG. 13

【図 14】

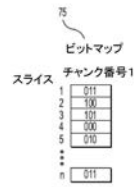


FIG. 14

【図 15】

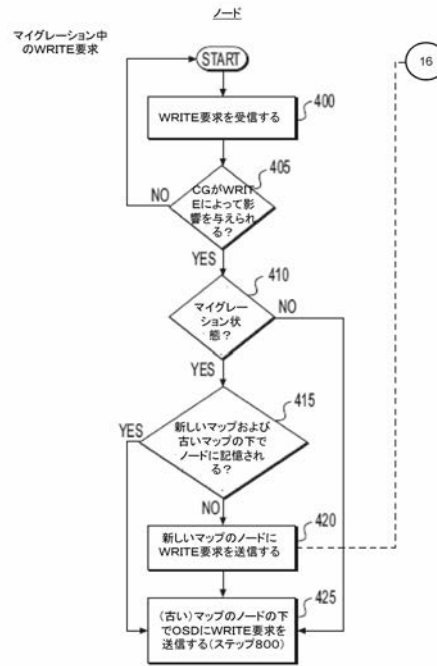


FIG. 15

【図 16】

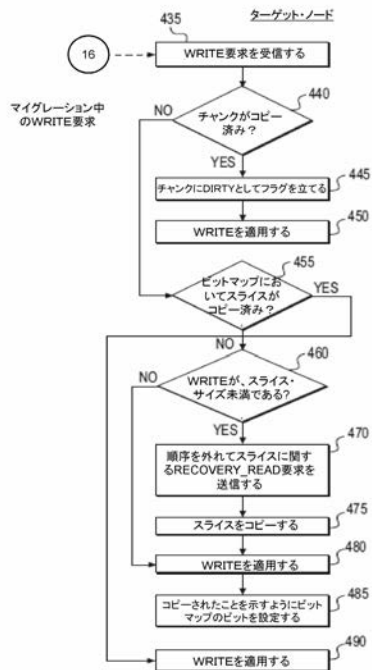


FIG. 16

【図 17】

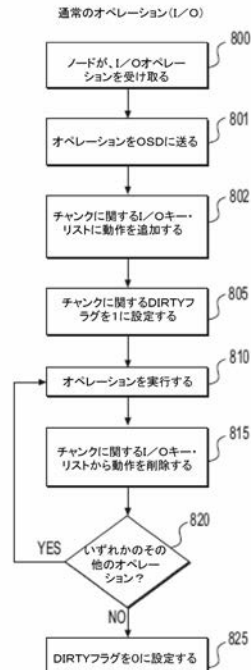


FIG. 17

【図 18】

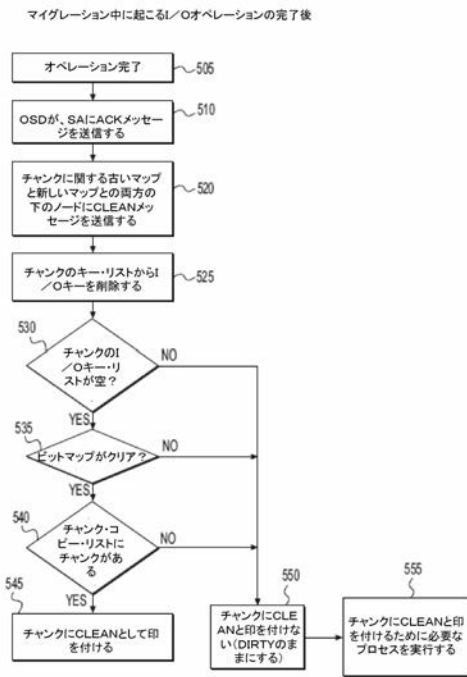


FIG. 18

【図 19】

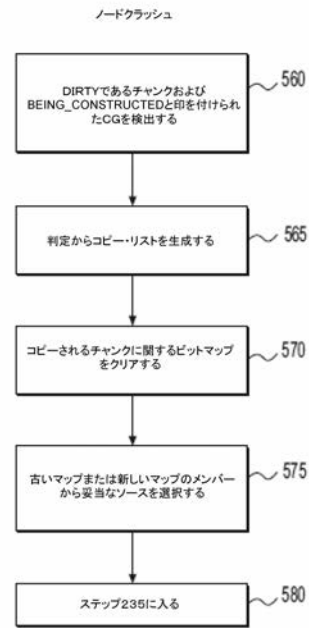


FIG. 19

【図 20】

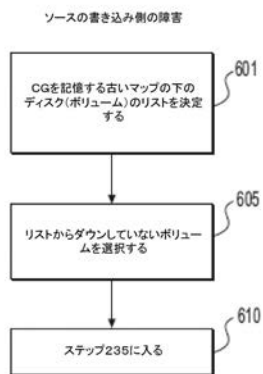


FIG. 20

【図 21】

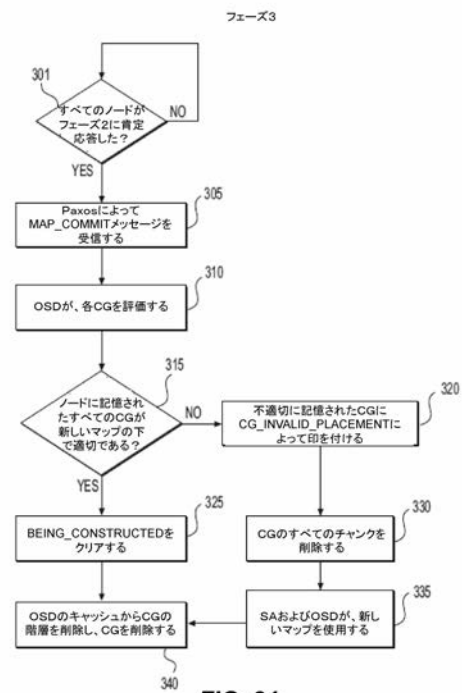


FIG. 21

【図 22】

ポリシー・テーブル

CG ID	ミラー・カウント
1	3
2	4
⋮	
n	

85

FIG. 22

【図 23】

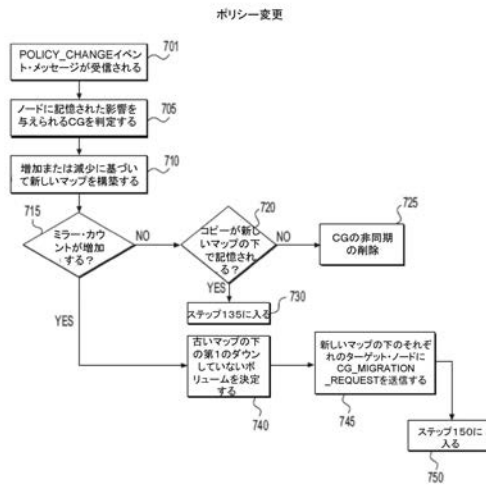


FIG. 23

【図 24】

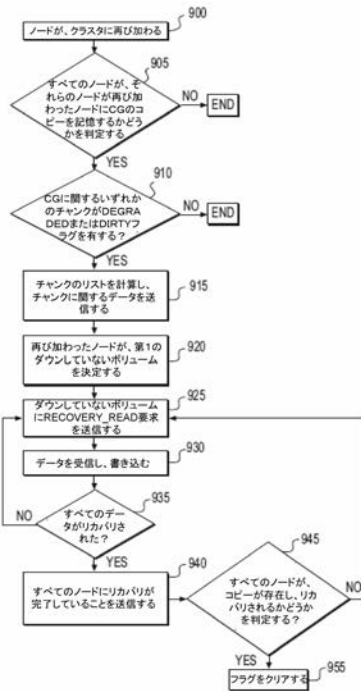


FIG. 24

【図 25】

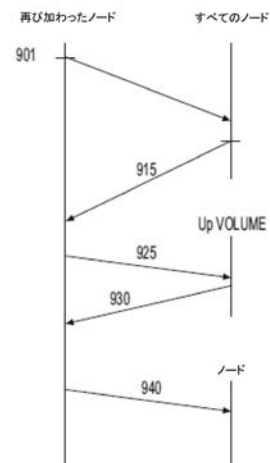


FIG. 25

【図 26】

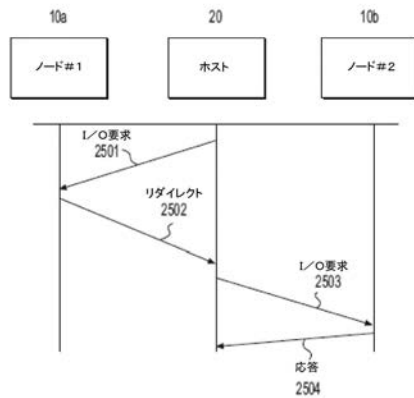


FIG. 26

【図 27】

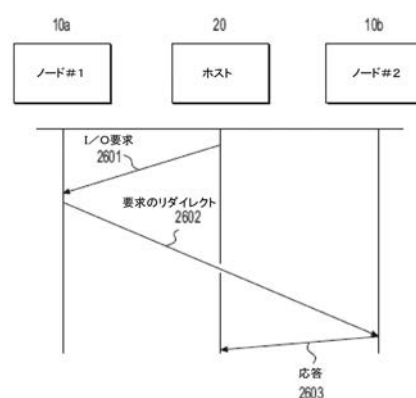


FIG. 27

【手続補正書】

【提出日】平成28年3月2日(2016.3.2)

【手続補正1】

【補正対象書類名】特許請求の範囲

【補正対象項目名】全文

【補正方法】変更

【補正の内容】

【特許請求の範囲】

【請求項1】

1以上のネットワークに接続された複数のノードを備え、

各ノードが、I/O(Input/Output)オペレーションを送受信可能であり、

各ノードが、データ及びデータレプリカを記憶するための1以上のストレージ・デバイスを備え、

各ノードが、データ及びデータレプリカの位置を示す位置情報に基づいて、前記複数のノードのストレージ・デバイスに記憶されたデータ及びデータレプリカの位置を特定できるようになっており、

各ノードと各ノードの各ストレージ・デバイスとの状態を示す状態情報が、

前記複数のノードのそれぞれに提供され、

ノード又はそのノードのストレージ・デバイスがアクティブ化されるか又は非アクティブ化されると、そのアクティブ化又は非アクティブ化に基づいて、変更され、

前記アクティブ化が、新しくアクティブ化されたノード又は新しくアクティブ化されたストレージ・デバイスを有するノードを含む場合、各ノードが、前記状態情報の変更に基づいて、前記複数のノードのストレージ・デバイスに記憶されたデータ又はデータレプリカの整合性を維持するためにリカバリオペレーションが必要とされるかどうかを判定し、

必要があれば、各ノードが、影響を与えられるデータを記憶する他ノードとの間でメッ

セージを送受信することによって、そのノードのストレージ・デバイスに記憶されたデータを復元するために、他ノードとは独立して前記リカバリオペレーションを実行する、ストレージ・システム。

【請求項 2】

データが、複数のノードにデータレプリカ又は消失訂正符号化セグメントとして冗長的に記憶され、

データレプリカ又は消失訂正符号化セグメントの数が、前記データの冗長性のポリシーに依存し、

データレプリカ又は消失訂正符号化セグメントの符号化が、前記データの前記冗長性のポリシー及び消失訂正符号化アルゴリズムに依存する

請求項 1 に記載のストレージ・システム。

【請求項 3】

各ノードが、イベントの通知を受信するようになっており、

前記イベントが、ノード又はそのノードのストレージ・デバイスのアクティブ化又は非アクティブ化である

請求項 2 に記載のストレージ・システム。

【請求項 4】

各イベントが、所定の順番で各ノードに届けられ、それら 1 以上のイベントが、一意で整合性のある番号である昇順のイベント番号を振られ、

各ノードが、前記 1 以上のイベントのそれぞれを受信するようになっており、

非アクティブ化の後でアクティブ化された各ノードが、前記ノードがアクティブ化されるときに前記イベントのそれぞれを受信する、

請求項 3 に記載のストレージ・システム。

【請求項 5】

前記状態情報が、

各ノードのストレージ・デバイスに記憶されたデータと各ノードのストレージ・デバイス上のデータレプリカ又は消失訂正符号化セグメントとの状態、及び、

前記他ノードのそれぞれのストレージ・デバイス上のデータとデータレプリカ又は消失訂正符号化セグメントとの状態

を表す情報を含み、

データ、データレプリカ、又は消失訂正符号化セグメントの前記状態が、

データを修正又は書き込みする前に、データ、データレプリカ、又は消失訂正符号セグメントに関して設定される D I R T Y フラグ、

データの修正又は書き込みがデータレプリカ又は消失訂正符号化セグメントの一部で成功したことを前記ノードが知っているときに、データ、データレプリカ、又は消失訂正符号化セグメントに関して設定される D E G R A D E D フラグ、

イベント後にデータが書き込まれたそのイベントの番号を示す情報、及び

イベント後に D E G R A D E D フラグが設定されたそのイベントの番号を示す情報、を含み、

前記 D I R T Y フラグ及び前記 D E G R A D E D フラグは、いずれも永続的であり、

前記 D I R T Y フラグが、前記データのデータレプリカ又は消失訂正符号化セグメントの少なくとも予め決められた最小限の数に対する書き込みが正常に完了したときにクリアされ、

前記データと前記データレプリカ又は前記消失訂正符号化セグメントとの位置が、永続的に記憶されることと、前記データの位置情報及びメタデータを基に算出されることのうちのいずれかが可能であり、

前記メタデータが、一意の I D、前記データの名前空間内の一意の名前属性、及び、前記データの内容の記述、のうちの少なくとも 1 つを含む

請求項 3 又は 4 に記載のストレージ・システム。

【請求項 6】

各ノードによって受信される前記イベントが、非アクティブ化の後にアクティブ化された第1ノードのアクティブ化又は前記第1ノードのストレージ・デバイスのアクティブ化であり、

前記判定が、前記位置情報に基づいて前記ノードのうちの1つが前記第1ノードに記憶されたデータ、データレプリカ、又は消失訂正符号化セグメントを有するかどうかに基づき、

前記第1ノードが、前記DEGRADEDフラグと、前記DIRTYフラグと、前記データの最後の修正又は書き込みのイベント番号を示す情報と、イベント後に前記DEGRADEDフラグが設定されたそのイベントの番号を示す情報とを含む情報のリストを受信し、

前記リストに基づいて、前記第1ノードが、どのデータ、データレプリカ、又は消失訂正符号化セグメントが再構築される必要があるかを判定し、

前記状態情報に基づいて、前記第1ノードが、前記第1ノードのストレージ・デバイス上のデータ、データレプリカ、又は消失訂正符号化セグメントを再構築するために再構築される必要がある前記第1ノード上のデータのデータレプリカ又は前記消失訂正符号化セグメントをどのノードが記憶するかを判定し、

前記第1ノードが、判定されたノードからの再構築されるデータ、データレプリカ、又は消失訂正符号化セグメントの内容を要求し、再構築すべきデータの内容を書き込み、完了すると、前記リストを送信した前記複数のノードのそれぞれにメッセージを送信し、

前記データのデータレプリカ又は消失訂正符号化セグメントのそれぞれが各ノードでリカバリされると判定される場合に、前記複数のノードのそれぞれが、前記データと、前記データレプリカと、前記消失訂正符号化セグメントとのいずれかの前記DEGRADEDフラグをクリアする、

請求項5に記載のストレージ・システム。

【請求項7】

請求項1乃至6のうちのいずれか1項に記載のストレージ・システムにおけるノードであって、

データ及びデータレプリカを記憶するための1以上のストレージ・デバイスと、

前記1以上のストレージ・デバイスに接続され、I/O (Input/Output) オペレーションを送受信でき、データ及びデータレプリカの位置を示す位置情報に基づいてデータ及びデータレプリカの位置を特定できる制御手段と

を備え、

前記ストレージ・システムの各ノードと各ノードの各ストレージ・デバイスとの状態を示す状態情報が、

前記複数のノードのそれぞれに提供され、

ノード又はそのノードのストレージ・デバイスがアクティブ化されるか又は非アクティブ化されると、そのアクティブ化又は非アクティブ化に基づいて、変更され、

前記アクティブ化が、新しくアクティブ化されたノード又は新しくアクティブ化されたストレージ・デバイスを有するノードを含む場合、そのノードの前記制御手段が、前記状態情報の変更に基づいて、前記複数のノードのストレージ・デバイスに記憶されたデータ又はデータレプリカの整合性を維持するためにリカバリオペレーションが必要とされるかどうかを判定し、

必要があれば、前記制御手段が、影響を与えられるデータを記憶する他ノードとの間でメッセージを送受信することによって、そのノードのストレージ・デバイスに記憶されたデータを復元するために、他ノードとは独立して前記リカバリオペレーションを実行する、

【請求項8】

I/O (Input/Output) オペレーションを送受信できデータ及びデータレプリカの位置を示す位置情報に基づいてデータ及びデータレプリカの位置を特定できる各ノードと、各

ノードの各ストレージ・デバイスとの状態を示す状態情報を、複数のノードのそれぞれに提供するステップと、

ノード又はそのノードのストレージ・デバイスがアクティブ化されるか又は非アクティブ化されると、そのアクティブ化又は非アクティブ化に基づいて、前記状態情報を変更するステップと、

前記アクティブ化が、新しくアクティブ化されたノード又は新しくアクティブ化されたストレージ・デバイスを有するノードを含む場合、前記状態情報の変更に基づいて、前記複数のノードのストレージ・デバイスに記憶されたデータ又はデータレプリカの整合性を維持するためにリカバリオペレーションが必要とされるかどうかを判定するステップと、

必要があれば、影響を与えられるデータを記憶する他ノードとの間でメッセージを送受信することによって、ストレージ・デバイスに記憶されたデータを復元するために、他ノードとは独立して前記リカバリオペレーションを実行するステップと
をコンピュータに実行させるコンピュータ・プログラム。

【請求項 9】

データが、複数のノードにデータレプリカ又は消失訂正符号化セグメントとして冗長的に記憶され、

データレプリカ又は消失訂正符号化セグメントの数が、前記データの冗長性のポリシーに依存し、

データレプリカ又は消失訂正符号化セグメントの符号化が、前記データの前記冗長性のポリシー及び消失訂正符号化アルゴリズムに依存する

請求項 8 に記載のコンピュータ・プログラム。

【請求項 10】

各ノードが、イベントの通知を受信するようになっており、

前記イベントが、ノード又はそのノードのストレージ・デバイスのアクティブ化又は非アクティブ化である

請求項 9 に記載のコンピュータ・プログラム。

【請求項 11】

各イベントが、所定の順番で各ノードに届けられ、それら 1 以上のイベントが、一意で整合性のある番号である昇順のイベント番号を振られ、

各ノードが、前記イベントのそれぞれを受信するようになっており、

非アクティブ化の後でアクティブ化された各ノードが、前記ノードがアクティブ化されるときに前記イベントのそれぞれを受信する

請求項 10 に記載のコンピュータ・プログラム。

【請求項 12】

前記状態情報が、

各ノードのストレージ・デバイスに記憶されたデータと各ノードのストレージ・デバイス上のデータレプリカ又は消失訂正符号化セグメントとの状態、及び、

前記他ノードのそれぞれのストレージ・デバイス上のデータとデータレプリカ又は消失訂正符号化セグメントとの状態

を表す情報を含み、

データ、データレプリカ、又は消失訂正符号化セグメントの前記状態が、

データを修正又は書き込みする前に、データ、データレプリカ、又は消失訂正符号セグメントに関して設定される D I R T Y フラグ、

データの修正又は書き込みがデータレプリカ又は消失訂正符号化セグメントの一部で成功したことを前記ノードが知っているときに、データ、データレプリカ、又は消失訂正符号化セグメントに関して設定される D E G R A D E D フラグ、

イベント後にデータが書き込まれたそのイベントの番号を示す情報、及び

イベント後に D E G R A D E D フラグが設定されたそのイベントの番号を示す情報、
を含み、

前記 D I R T Y フラグ及び前記 D E G R A D E D フラグは、いずれも永続的であり、

前記 D I R T Y フラグが、前記データのデータレプリカ又は消失訂正符号化セグメントの少なくとも予め決められた最小限の数に対する書き込みが正常に完了したときにクリアされ、

前記データと前記データレプリカ又は前記消失訂正符号化セグメントとの位置が、永続的に記憶されることと、前記データの位置情報及びメタデータを基に算出されることのうちのいずれかが可能であり、

前記メタデータが、一意の I D、前記データの名前空間内の一意の名前属性、及び、前記データの内容の記述、のうちの少なくとも 1 つを含む
請求項 10 又は 11 に記載のコンピュータ・プログラム。

【請求項 13】

各ノードによって受信される前記イベントが、非アクティブ化の後にアクティブ化された第 1 ノードのアクティブ化又は前記第 1 ノードのストレージ・デバイスのアクティブ化であり、

前記判定が、前記位置情報に基づいて前記ノードのうちの 1 つが前記第 1 ノードに記憶されたデータ、データレプリカ、又は消失訂正符号化セグメントを有するかどうかに基づき、

前記第 1 ノードが、前記 D E G R A D E D フラグと、前記 D I R T Y フラグと、前記データの最後の修正又は書き込みのイベント番号を示す情報と、イベント後に前記 D E G R A D E D フラグが設定されたそのイベントの番号を示す情報とを含む情報のリストを受信し、

前記リストに基づいて、前記第 1 ノードが、どのデータ、データレプリカ、又は消失訂正符号化セグメントが再構築される必要があるかを判定し、

前記状態情報に基づいて、前記第 1 ノードが、前記第 1 ノードのストレージ・デバイス上のデータ、データレプリカ、又は消失訂正符号化セグメントを再構築するために再構築される必要がある前記第 1 ノード上のデータのデータレプリカ又は前記消失訂正符号化セグメントをどのノードが記憶するかを判定し、

前記第 1 ノードが、判定されたノードからの再構築されるデータ、データレプリカ、又は消失訂正符号化セグメントの内容を要求し、再構築すべきデータの内容を書き込み、完了すると、前記リストを送信した前記複数のノードのそれぞれにメッセージを送信し、

前記データのデータレプリカ又は消失訂正符号化セグメントのそれぞれが各ノードでリカバリされると判定される場合に、前記複数のノードのそれぞれが、前記データと、前記データレプリカと、前記消失訂正符号化セグメントとのいずれかの前記 D E G R A D E D フラグをクリアする、

請求項 12 に記載のコンピュータ・プログラム製品。

【請求項 14】

I / O (Input/Output) オペレーションを送受信できデータ及びデータレプリカの位置を示す位置情報に基づいてデータ及びデータレプリカの位置を特定できる各ノードと、各ノードの各ストレージ・デバイスとの状態を示す状態情報を、複数のノードのそれぞれに提供するステップと、

ノード又はそのノードのストレージ・デバイスがアクティブ化されるか又は非アクティブ化されると、そのアクティブ化又は非アクティブ化に基づいて、前記状態情報を変更するステップと、

前記アクティブ化が、新しくアクティブ化されたノード又は新しくアクティブ化されたストレージ・デバイスを有するノードを含む場合、前記状態情報の変更に基づいて、前記複数のノードのストレージ・デバイスに記憶されたデータ又はデータレプリカの整合性を維持するためにリカバリオペレーションが必要とされるかどうかを判定するステップと、

必要があれば、影響を与えられるデータを記憶する他ノードとの間でメッセージを送受信することによって、ストレージ・デバイスに記憶されたデータを復元するために、他ノードとは独立して前記リカバリオペレーションを実行するステップと
を有する方法。

【請求項 15】

請求項 8 乃至 13 のうちのいずれか 1 項に記載のコンピュータ・プログラムに基づいて実行される請求項 14 に記載の方法。

【 国際調査報告 】

INTERNATIONAL SEARCH REPORT

International application No

PCT/US2013/065623

A. CLASSIFICATION OF SUBJECT MATTER

INV. G06F11/14 G06F11/20
ADD.

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

G06F

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

EPO-Internal, WPI Data, IBM-TDB, COMPENDEX, INSPEC

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	US 8 401 997 B1 (TAWRI DEEPAK [US] ET AL) 19 March 2013 (2013-03-19) column 2, line 7 - line 67 column 3, line 55 - column 6, line 41 column 10, line 29 - column 12, line 24 column 16, line 60 - column 21, line 24 figures 1,4-6,8 -----	1-20
A	US 6 907 507 B1 (KISELEV OLEG [US] ET AL) 14 June 2005 (2005-06-14) column 1, line 7 - column 4, line 58 ----- -/-	1-6

☒ Further documents are listed in the continuation of Box C.☒ See patent family annex.

* Special categories of cited documents :

"A" document defining the general state of the art which is not considered to be of particular relevance

"E" earlier application or patent but published on or after the international filing date

"L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

"O" document referring to an oral disclosure, use, exhibition or other means

"P" document published prior to the international filing date but later than the priority date claimed

"T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

"X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

"Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

"&" document member of the same patent family

Date of the actual completion of the international search

14 January 2015

Date of mailing of the international search report

29/01/2015

Name and mailing address of the ISA/

European Patent Office, P.B. 5818 Patentlaan 2
NL - 2280 HV Rijswijk
Tel: (+31-70) 340-2040,
Fax: (+31-70) 340-3016

Authorized officer

Johansson, Ulf

INTERNATIONAL SEARCH REPORT

International application No

PCT/US2013/065623

C(Continuation). DOCUMENTS CONSIDERED TO BE RELEVANT		
Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	US 7 617 259 B1 (MUTH JOHN A [US] ET AL) 10 November 2009 (2009-11-10) column 1, line 24 - column 2, line 53 column 3, line 48 - column 4, line 29 column 5, line 11 - line 40 column 6, line 58 - column 12, line 58 column 13, line 48 - column 15, line 13 figures 1,2a,3-8	7-14
A	----- SAGE A WEIL ET AL: "CRUSH: Controlled, Scalable, Decentralized Placement of Replicated Data", SUPERCOMPUTING, 2006. SC '06. PROCEEDINGS OF THE ACM/IEEE SC 2006 CONFERENCE, IEEE, PI, 1 November 2006 (2006-11-01), pages 31-31, XP031044223, ISBN: 978-0-7695-2700-0 page 1, line 1 - page 4, line 1 page 5, line 14 - page 6, line 21 page 10, line 1 - page 11, line 3 -----	15-20

INTERNATIONAL SEARCH REPORT

International application No.
PCT/US2013/065623**Box No. II Observations where certain claims were found unsearchable (Continuation of Item 2 of first sheet)**

This international search report has not been established in respect of certain claims under Article 17(2)(a) for the following reasons:

1. ☐ Claims Nos.:
because they relate to subject matter not required to be searched by this Authority, namely:

2. ☐ Claims Nos.:
because they relate to parts of the international application that do not comply with the prescribed requirements to such an extent that no meaningful international search can be carried out, specifically:

3. ☐ Claims Nos.:
because they are dependent claims and are not drafted in accordance with the second and third sentences of Rule 6.4(a).

Box No. III Observations where unity of invention is lacking (Continuation of Item 3 of first sheet)

This International Searching Authority found multiple inventions in this international application, as follows:

see additional sheet

1. ☒ As all required additional search fees were timely paid by the applicant, this international search report covers all searchable claims.

2. ☐ As all searchable claims could be searched without effort justifying an additional fees, this Authority did not invite payment of additional fees.

3. ☐ As only some of the required additional search fees were timely paid by the applicant, this international search report covers only those claims for which fees were paid, specifically claims Nos.:

4. ☐ No required additional search fees were timely paid by the applicant. Consequently, this international search report is restricted to the invention first mentioned in the claims; it is covered by claims Nos.:

Remark on Protest

- ☐ The additional search fees were accompanied by the applicant's protest and, where applicable, the payment of a protest fee.
- ☐ The additional search fees were accompanied by the applicant's protest but the applicable protest fee was not paid within the time limit specified in the invitation.
- ☒ No protest accompanied the payment of additional search fees.

International Application No. PCT/ US2013/ 065623

FURTHER INFORMATION CONTINUED FROM PCT/ISA/ 210

This International Searching Authority found multiple (groups of) inventions in this international application, as follows:

1. claims: 1-6

Distributed shared-nothing storage system, comprising: a plurality of independent nodes in a cluster connected to one or more networks, each comprising one or more devices for storing data; each independent node is configured to send and receive input/output (I/O) operations; the data stored on the plurality of nodes in the cluster includes replicas of the data, each node is aware of the location of the data and data replicas stored on the storage devices in the cluster based on a location map; and state information, which indicates the status of each node and each storage device on each node, which is available to all nodes in the cluster, upon activation or inactivation of a node or activation or inactivation of a storage device on any node in the cluster, the state information is updated and changed based on the corresponding activation or inactivation, based on the change in the state information, each node in the cluster, including the newly activated node or node that has a newly activated storage device, determines if it must perform a recovery operation to maintain the consistency of any data or data replicas stored on the nodes in the cluster, and each node performs the recovery operation, if needed, independently from the other nodes in the cluster to restore the data stored on its storage devices by sending and receiving messages to other nodes in the cluster that store affected data.

2. claims: 7-14

Distributed shared-nothing storage system, comprising: a plurality of independent nodes in a cluster connected to a network, each comprising one or more storage devices for storing data; each node is configured to send and receive input/output (I/O) operations; and each node has memory for conducting local processing and is separate from the storage devices, wherein I/O operations are performed on individual nodes and each node that receives an I/O operation performs the I/O operation only for the data stored on its storage devices, wherein a location of data stored on the storage devices in the cluster is derived from layout information available to all the nodes, wherein the location of the data is stored and calculated by each node, as needed, from the layout information, and wherein, upon receiving a first event message indicating the occurrence of an event, each node evaluates all data stored on its storage devices and determines if any data is affected by information in the first event message.

3. claims: 15-20

International Application No. PCT/ US2013/ 065623

FURTHER INFORMATION CONTINUED FROM PCT/ISA/ 210

Distributed shared-nothing storage system, comprising: a plurality of independent nodes in a cluster connected to a network, each comprising one or more storage devices for storing data; each node is configured to send and receive input/output (I/O) operations; and state information available to all nodes which indicates whether each node and each storage device on each node is in an active state, wherein, data in the cluster is redundantly stored as replicas or erasure-coded segments on a plurality of nodes in the cluster, wherein the number of replicas or erasure-coded segments depends on the redundancy policy of the data, wherein the encoding of replicas or erasure-coded segments depends on the redundancy policy and erasure-coding algorithm of the data, wherein a location of the data and replicas or erasure-coded segments of the data on the storage devices of the nodes in the cluster is derived from layout information available to all the nodes, wherein a location of the data is stored and calculated as needed from the layout information, and wherein a redundancy policy for the storage system sets a mirror count number indicating the number of data replicas or erasure-coded segments of the data, data replicas or erasure-coded segments are migrated from a first location to a second location on the storage devices in the cluster or data, data replicas, or erasure-coded segments are removed from the cluster based on a change in the redundancy policy of the storage system.

INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No

PCT/US2013/065623

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
US 8401997	B1	19-03-2013	NONE
US 6907507	B1	14-06-2005	US 6907507 B1 14-06-2005
		US 7089385 B1	08-08-2006
US 7617259	B1	10-11-2009	NONE

フロントページの続き

(51)Int.Cl.

F I

テーマコード(参考)

G 0 6 F 12/00 5 3 1 R

(81)指定国 AP(BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, SZ, TZ, UG, ZM, ZW), EA(AM, AZ, BY, KG, KZ, RU, TJ, TM), EP(AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OA(BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG), AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JP, KE, KG, KN, KP, KR, KZ, LA, LC, LK, LR, LS, LT, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US

(72)発明者 キセレヴ, オレグ

アメリカ合衆国 9 4 3 0 1 カリフォルニア州, パロ アルト, ブライアント ストリート #
3 3 3 5 5 5

(72)発明者 ボール, ガウラブ

アメリカ合衆国 9 5 0 1 4 カリフォルニア州, クパチーノ, アpartment ディー, サウス
ブラニー アベニュー 1 0 1 0 9

(72)発明者 ヤングワース, クリストファー

アメリカ合衆国 9 5 1 2 0 カリフォルニア州, サンノゼ, ガダラジャラ コート 1 6 0 4

F ターム(参考) 5B027 CC04