

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号

特許第6983154号
(P6983154)

(45) 発行日 令和3年12月17日 (2021. 12. 17)

(24) 登録日 令和3年11月25日 (2021. 11. 25)

(51) Int. Cl. F I
GO 6 N 3/10 (2006. 01) GO 6 N 3/10
GO 6 F 9/50 (2006. 01) GO 6 F 9/50 1 5 O E

請求項の数 19 (全 20 頁)

(21) 出願番号	特願2018-521825 (P2018-521825)	(73) 特許権者	502208397
(86) (22) 出願日	平成28年10月28日 (2016. 10. 28)		グーグル エルエルシー
(65) 公表番号	特表2018-538607 (P2018-538607A)		G o o g l e L L C
(43) 公表日	平成30年12月27日 (2018. 12. 27)		アメリカ合衆国 カリフォルニア州 9 4
(86) 国際出願番号	PCT/US2016/059449		0 4 3 マウンテン ビュー アンフィシ
(87) 国際公開番号	W02017/075438		アター パークウェイ 1 6 0 0
(87) 国際公開日	平成29年5月4日 (2017. 5. 4)		1 6 0 0 Amphitheatre P
審査請求日	平成30年6月6日 (2018. 6. 6)		arkway 9 4 0 4 3 Mounta
審査番号	不服2020-6550 (P2020-6550/J1)		in View, CA U. S. A.
審査請求日	令和2年5月14日 (2020. 5. 14)	(74) 代理人	100108453
(31) 優先権主張番号	62/247, 709		弁理士 村山 靖彦
(32) 優先日	平成27年10月28日 (2015. 10. 28)	(74) 代理人	100110364
(33) 優先権主張国・地域又は機関	米国 (US)		弁理士 実広 信哉
		(74) 代理人	100133400
			弁理士 阿部 達彦

最終頁に続く

(54) 【発明の名称】 計算グラフの処理

(57) 【特許請求の範囲】

【請求項 1】

計算グラフを処理するコンピュータ実施方法であって、

計算グラフを処理する要求をクライアントから受信するステップと、

前記計算グラフを表すデータを取得するステップであって、前記計算グラフは、複数のノードおよび有向エッジを含み、各ノードは、それぞれの演算を表し、各有向エッジは、それぞれの第1のノードによって表される演算の出力を入力として受信する演算を表すそれぞれの第2のノードに前記それぞれの第1のノードを接続する、ステップと、

前記計算グラフを処理するための複数の利用可能なデバイスを選択するステップと、

前記複数の利用可能なデバイスを使用して前記計算グラフを処理するステップであって

10

、

ノードのグループに向けて有向エッジ上を流れる共有データを演算する前記ノードのグループを特定するために前記計算グラフを分析するステップと、

複数のサブグラフに前記計算グラフを分割するステップであって、各サブグラフは、前記計算グラフ内の1つまたは複数のノードを含み、前記分割するステップは、特定されたグループの各々について、ノードの前記特定されたグループを含むそれぞれのサブグラフを生成するステップを含み、前記複数のサブグラフは、入力として同一のデータを受信する複数のノードを含むサブグラフを含む、ステップと、

各サブグラフについて、前記サブグラフ内の前記1つまたは複数のノードによって表される前記演算を、処理するための前記複数の利用可能なデバイスのうちのそれぞれの利

20

用可能なデバイスに、前記それぞれの利用可能なデバイスの利用可能なメモリサイズに基づいて割り当てるステップであって、前記計算グラフの前記1つまたは複数のノードによって表される前記演算は、ニューラルネットワークのための推論または訓練演算である、ステップと、

前記デバイスに割り当てられた前記演算を各デバイスに行わせることによって前記計算グラフを処理するステップと

を含む、ステップと

を含む、方法。

【請求項 2】

前記要求は、1つまたは複数のそれぞれのノードから1つまたは複数の特定の出力を指定し、

前記1つまたは複数のそれぞれのノードが割り当てられたデバイスから、前記1つまたは複数の特定の出力を受信するステップと、

前記1つまたは複数の特定の出力を前記クライアントに提供するステップと

をさらに含む、請求項1に記載の方法。

【請求項 3】

前記要求は、複数の所定のサブグラフに前記計算グラフを分割するラベルを含み、前記計算グラフを分割するステップは、前記複数の所定のサブグラフに前記計算グラフを分割するステップを含む、請求項1または2に記載の方法。

【請求項 4】

前記複数の利用可能なデバイスのうちの各デバイスは、前記複数の利用可能なデバイスのうちの他のデバイスから独立した演算を行うハードウェアリソースである、請求項1から3のいずれか一項に記載の方法。

【請求項 5】

各サブグラフについて、前記サブグラフ内の前記1つまたは複数のノードによって表される前記演算を、それぞれのデバイスに割り当てるステップは、前記演算を、前記サブグラフ内の前記ノードによって表される前記演算を行うために必要な計算性能を有するデバイスに割り当てるステップを含む、請求項1から4のいずれか一項に記載の方法。

【請求項 6】

チェーン構造内に配置されているノードのグループを特定するために前記計算グラフを分析するステップをさらに含み、

前記分割するステップは、特定されたグループの各々について、ノードの前記特定されたグループを含むそれぞれのサブグラフを生成するステップを含む、請求項1から5のいずれか一項に記載の方法。

【請求項 7】

デバイスに対するサブグラフ内の1つまたは複数のノードによって表される演算の初期割り当てを決定するステップと、

統計値を決定するために前記デバイスをモニタするステップと、

前記統計値を使用して前記初期割り当てを調整するステップと、

前記調整した初期割り当てに基づいて前記サブグラフの前記演算を前記デバイスに再割り当てするステップと

をさらに含む、請求項1から6のいずれか一項に記載の方法。

【請求項 8】

改善についての閾値量に達するまで、前記モニタするステップと、前記調整するステップと、前記再割り当てするステップとを繰り返すステップをさらに含む、請求項7に記載の方法。

【請求項 9】

前記統計値は、各サブグラフについてのそれぞれの動作時間またはそれぞれのアイドル時間を含む、請求項7に記載の方法。

【請求項 10】

10

20

30

40

50

モデル入力を受信するステップと、処理済み計算グラフによって表される演算に従って前記モデル入力を処理するステップとをさらに含む、請求項1から9のいずれか一項に記載の方法。

【請求項 1 1】

請求項1から9のいずれか一項に記載の方法によって取得された処理済み計算グラフに対応する機械学習モデルを提供するステップと、モデル入力を前記機械学習モデルを使用して処理するステップとを含む、方法。

【請求項 1 2】

システムであって、

1つまたは複数のコンピュータと、

10

前記1つまたは複数のコンピュータに接続されるとともに命令を記憶するコンピュータ可読媒体とを含み、前記命令は、前記1つまたは複数のコンピュータによって実行されると、前記1つまたは複数のコンピュータに、ニューラルネットワーク層の各々について、計算グラフを割り当てる要求をクライアントから受信するステップと、

前記計算グラフを表すデータを取得するステップであって、前記計算グラフは、複数のノードおよび有向エッジを含み、各ノードは、それぞれの演算を表し、各有向エッジは、それぞれの第1のノードによって表される演算の出力を入力として受信する演算を表すそれぞれの第2のノードに前記それぞれの第1のノードを接続する、ステップと、

前記計算グラフを処理するための複数の利用可能なデバイスを特定するステップと、

前記複数の利用可能なデバイスを使用して前記計算グラフを処理するステップであって

20

、
ノードのグループに向けて有向エッジ上を流れる共有データを演算する前記ノードのグループを特定するために前記計算グラフを分析するステップと、

複数のサブグラフに前記計算グラフを分割するステップであって、各サブグラフは、前記計算グラフ内の1つまたは複数のノードを含み、前記分割するステップは、特定されたグループの各々について、ノードの前記特定されたグループを含むそれぞれのサブグラフを生成するステップを含み、前記複数のサブグラフは、入力として同一のデータを受信する複数のノードを含むサブグラフを含む、ステップと、

各サブグラフについて、前記サブグラフ内の前記1つまたは複数のノードによって表される前記演算を、処理するための前記複数の利用可能なデバイスのうちのそれぞれの利用可能なデバイスに、前記それぞれの利用可能なデバイスの利用可能なメモリサイズに基づいて割り当てるステップと、

30

前記デバイスに割り当てられた前記演算を各デバイスに行わせることによって前記計算グラフを処理するステップと

を含む、ステップと

を含む動作を行わせる、システム。

【請求項 1 3】

前記動作は、

デバイスに対するサブグラフ内の1つまたは複数のノードによって表される演算の初期割り当てを決定するステップと、

40

統計値を決定するために前記デバイスをモニタするステップと、

前記統計値を使用して前記初期割り当てを調整するステップと、

前記調整した初期割り当てに基づいて前記サブグラフの前記演算を前記デバイスに再割り当てするステップと

をさらに含む、請求項12に記載のシステム。

【請求項 1 4】

前記演算は、改善についての閾値量に達するまで、前記モニタするステップと、前記調整するステップと、前記再割り当てするステップとを繰り返すステップをさらに含む、請求項13に記載のシステム。

【請求項 1 5】

50

前記統計値は、各サブグラフについてのそれぞれの動作時間またはそれぞれのアイドル時間を含む、請求項13に記載のシステム。

【請求項16】

命令を有するコンピュータ可読媒体であって、前記命令は、1つまたは複数のコンピュータによって実行されると、前記1つまたは複数のコンピュータに、

計算グラフを割り当てる要求をクライアントから受信するステップと、

前記計算グラフを表すデータを取得するステップであって、前記計算グラフは、複数のノードおよび有向エッジを含み、各ノードは、それぞれの演算を表し、各有向エッジは、それぞれの第1のノードによって表される演算の出力を入力として受信する演算を表すそれぞれの第2のノードに前記それぞれの第1のノードを接続する、ステップと、

前記計算グラフを処理するための複数の利用可能なデバイスを特定するステップと、

前記複数の利用可能なデバイスを使用して前記計算グラフを処理するステップであって、

ノードのグループに向けて有向エッジ上を流れる共有データを演算する前記ノードのグループを特定するために前記計算グラフを分析するステップと、

複数のサブグラフに前記計算グラフを分割するステップであって、各サブグラフは、前記計算グラフ内の1つまたは複数のノードを含み、前記分割するステップは、特定されたグループの各々について、ノードの前記特定されたグループを含むそれぞれのサブグラフを生成するステップを含み、前記複数のサブグラフは、入力として同一のデータを受信する複数のノードを含むサブグラフを含む、ステップと、

各サブグラフについて、前記サブグラフ内の前記1つまたは複数のノードによって表される前記演算を、処理するための前記複数の利用可能なデバイスのうちのそれぞれの利用可能なデバイスに、前記それぞれの利用可能なデバイスの利用可能なメモリサイズに基づいて割り当てるステップであって、前記計算グラフの前記1つまたは複数のノードによって表される前記演算は、ニューラルネットワークのための推論または訓練演算である、ステップと、

前記デバイスに割り当てられた前記演算を各デバイスに行わせることによって前記計算グラフを処理するステップと

を含む、ステップと

を含む動作を行わせる、コンピュータ可読媒体。

【請求項17】

前記動作は、

デバイスに対するサブグラフ内のノードによって表される演算の初期割り当てを決定するステップと、

統計値を決定するために前記デバイスをモニタするステップと、

前記統計値を使用して前記初期割り当てを調整するステップと、

前記調整した初期割り当てに基づいて前記演算を前記デバイスに再割り当てするステップと

をさらに含む、請求項16に記載のコンピュータ可読媒体。

【請求項18】

前記動作は、改善についての閾値量に達するまで、前記モニタするステップと、前記調整するステップと、前記再割り当てするステップとを繰り返すステップをさらに含む、請求項17に記載のコンピュータ可読媒体。

【請求項19】

前記統計値は、各サブグラフについてのそれぞれの動作時間またはそれぞれのアイドル時間を含む、請求項17または18に記載のコンピュータ可読媒体。

【発明の詳細な説明】

【技術分野】

【0001】

本明細書は、ニューラルネットワークを表している計算グラフを処理することおよび/

10

20

30

40

50

またはモデル入力処理するための処理済み計算グラフの使用に関する。

【背景技術】

【0002】

ニューラルネットワークは、1つまたは複数の層のモデルを使用して、受信した入力に対する、例えば、1つまたは複数の分類といった、出力を生成する、機械学習モデルである。いくつかのニューラルネットワークは、出力層に加えて1つまたは複数の隠れ層を含む。各隠れ層の出力は、ネットワーク内の次段の層、すなわち、次段の隠れ層またはネットワークの出力層に対する入力として使用される。ネットワークの各層は、その層のためのそれぞれのパラメータのセットの現在の値に従って、受信した入力から出力を生成する。

10

【0003】

ニューラルネットワークの複数の層は、個々のデバイスによって処理され得る。デバイスは、例えば、入力からある層における出力を生成するといった、演算を行うプロセッサを有し得るし、演算による出力をメモリに記憶する。ニューラルネットワークにおいて出力を生成するのに必要とされる演算の数およびサイズが一般的には膨大であるため、1つのデバイスでは、ニューラルネットワークの複数の層を処理するのにかなりの量の時間がかかり得る。

【発明の概要】

【課題を解決するための手段】

【0004】

20

概して、本明細書は、ニューラルネットワークまたは別の機械学習モデルを表している計算グラフを処理するためのシステムおよび方法を説明している。

【0005】

一般に、本明細書において説明した発明特定事項の1つの革新的な態様は、計算グラフを処理する要求をクライアントから受信するステップと、計算グラフを表すデータを取得するステップであって、計算グラフは、複数のノードおよび有向エッジを含み、各ノードは、それぞれの演算を表し、各有向エッジは、それぞれの第1のノードによって表される演算の出力を入力として受信する演算を表すそれぞれの第2のノードにそれぞれの第1のノードを接続する、ステップと、要求された演算を行うための複数の利用可能なデバイスを特定するステップと、複数のサブグラフに計算グラフを分割するステップであって、各サブグラフは、計算グラフ内の1つまたは複数のノードを含む、ステップと、各サブグラフについて、サブグラフ内の1つまたは複数のノードによって表される演算を、演算のために複数の利用可能なデバイスのうちのそれぞれの利用可能なデバイスに割り当てるステップとのアクションを含む方法で具現化され得る。方法は、コンピュータ実施方法であってもよい。

30

【0006】

実施形態は、以下の特徴のうちの1つまたは複数を含み得る。要求は、1つまたは複数のそれぞれのノードから1つまたは複数の特定の出力を指定し、1つまたは複数のそれぞれのノードが割り当てられたデバイスから、1つまたは複数の特定の出力を受信するステップと、1つまたは複数の特定の出力をクライアントに提供するステップをさらに含む。計算グラフによる演算は、ニューラルネットワークのための推論または訓練演算を含む。要求は、複数の所定のサブグラフに計算グラフを分割するラベルを含み、計算グラフを分割するステップは、複数の所定のサブグラフに計算グラフを分割するステップを含む。各デバイスは、複数のデバイスのうちの他のデバイスから独立した演算を行うハードウェアリソースである。各サブグラフをそれぞれのデバイスに割り当てるステップは、サブグラフを、サブグラフ内のノードによって表される演算を行うために必要な計算性能を有するデバイスに割り当てるステップを含む。チェーン構造内に配置されているノードのグループを特定するために計算グラフを分析するステップをさらに含む、分割するステップは、特定されたグループの各々について、ノードの特定されたグループを含むそれぞれのサブグラフを生成するステップを含む。ノードのグループに向けて有向エッジ上を流れる共有デー

40

50

タを演算するノードのグループを特定するために計算グラフを分析するステップをさらに含み、分割するステップは、特定されたグループの各々について、ノードの特定されたグループを含むそれぞれのサブグラフを生成するステップを含む。デバイスに対するサブグラフの初期割り当てを決定するステップと、統計値を決定するためにデバイスをモニタするステップと、統計値を使用して初期割り当てを調整するステップと、調整した初期割り当てに基づいてサブグラフをデバイスに再割り当てするステップとをさらに含む。改善についての閾値量に達するまで、モニタするステップと、調整するステップと、再割り当てするステップとを繰り返すステップをさらに含む。統計値は、各サブグラフについてのそれぞれの動作時間またはそれぞれのアイドル時間を含む。

【0007】

10

さらなる実施形態においては、方法は、モデル入力を受信するステップと、処理済み計算グラフによって表される演算に従ってモデル入力を処理するステップとをさらに含む。

【0008】

本明細書において説明した発明特定事項の別の革新的な態様は、第1の態様の方法によって取得された処理済み計算グラフに対応する機械学習モデルを提供するステップと、モデル入力を機械学習モデルを使用して処理するステップとのアクションを含み得る方法で具現化され得る。モデル入力の処理は、機械学習モデルを訓練することの一部であり得る、または、それは、モデル入力から推論を生成することの一部であり得る。

【0009】

別の態様においては、本明細書において説明した発明特定事項は、複数のデバイスによって、第1の態様の方法によって取得された処理済み計算グラフを実行するステップのアクションを含み得る方法で具現化され得る。

20

【0010】

これらの態様においては、計算グラフは、例えば、ニューラルネットワークなどといった、機械学習モデルの表現であり得る。

【0011】

本明細書において説明した発明特定事項の別の革新的な態様は、複数のデバイスを使用して計算グラフに従ってモデル入力を処理するアクションを含む方法であって、計算グラフは、複数のノードおよび有向エッジを含み、各ノードは、それぞれの演算を表し、各有向エッジは、それぞれの第1のノードによって表される演算の出力を入力として受信する演算を表すそれぞれの第2のノードにそれぞれの第1のノードを接続し、方法は、複数のデバイスの各々について、デバイスに割り当てられた計算グラフのサブグラフを表すデータを受信するステップであって、サブグラフは、複数のノードおよび計算グラフからの有向エッジを含む、ステップと、サブグラフ内のノードによって表される演算を行うステップとを含む、方法で具現化され得る。

30

【0012】

本態様の実施形態は、以下の特徴のうちの1つまたは複数を含み得る。要求は、1つまたは複数のそれぞれのノードから1つまたは複数の特定の出力を指定し、サブグラフ内の1つまたは複数のそれぞれのノードから1つまたは複数の特定の出力を特定する要求を受信するステップと、1つまたは複数の特定の出力をクライアントに提供するステップとをさらに含む。方法は、統計値をモニタするステップと、統計値をクライアントに提供するステップとをさらに含む。統計値は、各サブグラフについてのそれぞれの動作時間またはそれぞれのアイドル時間を含む。サブグラフ内のノードによって表される演算を行うステップは、非同期的に演算を行うステップを含む。非同期的に演算を行うステップは、キュー、非ブロッキングカーネル、またはその両方を使用して演算を行うステップを含む。

40

【0013】

他の態様は、上記態様のいずれか1つに対応するシステムおよびコンピュータ可読媒体を提供している。コンピュータ可読媒体は、非一時的コンピュータ可読媒体であり得るが、発明はこれに限定されない。

【0014】

50

本明細書において説明した発明特定事項の特定の実施形態が、以下の利点の1つまたは複数を実現するために実施され得る。ニューラルネットワークの演算、例えば、入力から推論を生成する演算といった演算は、ノードおよび有向エッジの計算グラフとして表すことができる。システムは、このような計算グラフ表現を処理して、ニューラルネットワークの演算を効率的に行う。システムは、計算グラフが一連の層として表される従来のニューラルネットワークと比べてより少ない抽象化を有することになるため、このような効率性を達成している。具体的には、計算グラフを、従来のニューラルネットワーク表現と比べて、並列演算のためにより容易に分割することが可能である。例として、計算グラフのサブグラフを、一意なデバイスに割り当てることができ、例えば、各サブグラフを、他のサブグラフとは異なるデバイスに割り当てることができ、それらの各々は、それぞれのサブグラフ内の演算を行って、ニューラルネットワークの演算を行うために必要な時間全体を低減している。

10

【0015】

本明細書の発明特定事項の1つまたは複数の実施形態についての詳細を以下の添付の図面および説明に記載している。発明特定事項の他の特徴、態様、および利点は、説明、図面、および特許請求の範囲から自明となるであろう。態様および実施形態を組み合わせ得ること、および1つの態様または実施形態に関連して説明した特徴を他の態様または実施形態に関連して実施し得ることは諒解されるであろう。

【図面の簡単な説明】

【0016】

20

【図1】計算グラフとして表されるニューラルネットワークのための演算を分散するための例示的な計算グラフシステムの図である。

【図2】計算グラフを処理するための例示的な方法のフロー図である。

【図3】例示的な計算グラフである。

【図4】サブグラフをデバイスに割り当てるための例示的なプロセスのフロー図である。

【発明を実施するための形態】

【0017】

様々な図面中の類似の参照番号および記号は類似の要素を示す。

【0018】

本明細書は、分散方式で計算グラフによって表される演算を行う計算グラフシステムを一般的に説明している。

30

【0019】

計算グラフは、有向エッジによって接続されるノードを含む。計算グラフ内の各ノードは演算を表す。ノードへの入力方向エッジは、ノードに対する入力、すなわち、ノードによって表される演算に対する入力のフローを表す。ノードからの出力方向エッジは、別のノードによって表される演算に対する入力として使用されることになる、そのノードによって表される演算の出力のフローを表す。そのため、グラフ内の第1のノードをグラフ内の第2のノードに接続する有向エッジは、第1のノードによって表される演算によって生成された出力が第2のノードによって表される演算に対する入力として使用されることを示している。

40

【0020】

一般的に、計算グラフ内の有向エッジに沿って流れる入力および出力がテンソルである。テンソルは、配列の次元に対応する固有の次数を有する、数値または例えば文字列といった他の値の多次元配列である。例えば、スカラー値は、0次テンソルであり、数値のベクトルは、1次テンソルであり、行列は、2次テンソルである。

【0021】

いくつかの実施形態においては、計算グラフ内に表されている演算は、ニューラルネットワーク演算または異なる種類の機械学習モデルのための演算である。ニューラルネットワークは、非線形ユニットの1つまたは複数の層を使用して受信した入力に対する出力を予測する機械学習モデルである。いくつかのニューラルネットワークは、出力層に加えて

50

1つまたは複数の隠れ層を含むディープニューラルネットワークである。各隠れ層の出力は、ネットワーク内の別の層、すなわち、別の隠れ層、出力層、またはその両方への入力として使用される。ネットワークのいくつかの層は、それぞれのパラメータのセットの現在の値に従って受信した入力から出力を生成する一方で、ネットワークの他の層は、パラメータを有し得ない。

【0022】

例えば、計算グラフによって表される演算は、ニューラルネットワークが推論を計算するために、すなわち、ニューラルネットワークの複数の層を介して入力を処理して入力に対するニューラルネットワーク出力を生成するために必要な演算であり得る。別の例としては、計算グラフによって表される演算は、ニューラルネットワークのパラメータの値を調整するために、例えば、パラメータの初期値からパラメータの訓練済みの値を決定するために、ニューラルネットワーク訓練プロセスを行うことによって、ニューラルネットワークを訓練するのに必要な演算であり得る。いくつかのケースにおいては、例えば、ニューラルネットワークの訓練中に、計算グラフによって表される演算は、ニューラルネットワークの複数のレプリカによって行われる演算を含み得る。

10

【0023】

例として、前段の層から入力を受信するあるニューラルネットワーク層は、パラメータ行列を使用してパラメータ行列と入力との間の行列乗算を行い得る。いくつかのケースにおいては、この行列乗算は、計算グラフ内の複数のノードとして表され得る。例えば、行列乗算は、複数の積和演算に分割され得るし、各演算は、計算グラフ内の異なるノードによって表され得る。各ノードによって表される演算は、後続のノードに向けて有向エッジ上を流れるそれぞれの出力を生成し得る。最終ノードによって表される演算が行列乗算の結果を生成した後に、結果は、別のノードに向けて有向エッジ上を流れる。結果は、行列乗算を行うニューラルネットワーク層の出力に等しい。

20

【0024】

いくつかの他のケースにおいては、行列乗算は、グラフ内の1つのノードとして表される。ノードによって表される演算は、入力として、第1の有向エッジに対する入力テンソルと、第2の有向エッジに対する、例えば、パラメータ行列といった、重みテンソルとを受信し得る。ノードは、第3の有向エッジ上で、処理、例えば、入力と重みテンソルとの行列乗算を行い、ニューラルネットワーク層の出力に等しい出力テンソルを出力し得る。

30

【0025】

計算グラフ内のノードによって表され得る他のニューラルネットワーク演算は、例えば、減算、除算、および勾配の計算といった、他の数学的演算と、例えば、連結、スプライス、分割、またはランク付けといった、配列演算とを含み、ニューラルネットワークは、例えば、ソフトマックス、シグモイド、正規化線形ユニット(ReLU)、または畳み込みといった、ブロック演算を構築する。

【0026】

特に、ニューラルネットワークのための演算を、異なるハードウェアプロファイルを有する複数のデバイスにわたって分散する場合には、ニューラルネットワークを計算グラフとして表すことは、ニューラルネットワークを効率的に実装する柔軟かつ粒度の細かい方法を提供することになる。

40

【0027】

図1は、計算グラフとして表されるニューラルネットワークのための演算を分散するための例示的な計算グラフシステム100を図示している。システム100は、以下で説明しているシステム、コンポーネント、および技法を実施し得る、1つまたは複数のロケーションにある1つまたは複数のコンピュータ上のコンピュータプログラムとして実施されるシステムの例である。

【0028】

クライアント102のユーザは、演算がニューラルネットワークを表す計算グラフ上で行われるように要求し得る。クライアント102は、コンピュータ上で動作しているアプリケ

50

ーションであってもよい。

【0029】

要求の一部として、クライアント102は、計算グラフを特定するデータをシステム100に提供し、計算グラフ上で行われることになる演算のタイプを指定する。

【0030】

例えば、要求は、特定のニューラルネットワークのための推論を表す計算グラフを特定し得るし、推論が行われるべき入力を特定し得る。

【0031】

別の例としては、要求は、特定のニューラルネットワークのための訓練プロシージャを表す計算グラフを特定し得るし、訓練が行われるべき訓練データなどの入力を特定し得る。この例においては、訓練プロシージャを表す計算グラフを処理する要求を受信すると、システム100は、例えば、従来の逆伝播技法または他のニューラルネットワーク訓練技法を使用して、計算グラフの1つまたは複数のエッジのためのパラメータについての修正された値を決定し得る。システム100は、修正されたパラメータをデバイスのメモリに記憶し得るし、実行器106は、読み出して、システム100において、修正された重みのアドレスを記憶し得る。修正された重みを必要とする推論、訓練、または他の演算のためのクライアント102からのさらなる要求には、システム100は、前記アドレスを使用して修正された重みにアクセスし得る。

【0032】

いくつかのケースにおいては、要求は、要求に応じて送信されるべき応答を指定し得る。例えば、ニューラルネットワーク訓練要求については、クライアント102は、要求されたニューラルネットワーク訓練演算を完了し終えたということについてのインディケーション、必要に応じて、ニューラルネットワークのパラメータの訓練済みの値またはクライアント102が訓練済みの値にアクセスすることができるメモリロケーションのインディケーションを要求し得る。別の例としては、ニューラルネットワーク推論要求については、クライアント102は、計算グラフの1つまたは複数の特定のノードからの推論演算を表す出力値を要求し得る。

【0033】

システム100は、計算グラフによって表される演算を複数のデバイス116~122にわたって分割することによって、特定の出力を生成する演算を行う。システム100は、データ通信ネットワーク114、例えば、ローカルエリアネットワーク(LAN)またはワイドエリアネットワーク(WAN)の先にある複数のデバイス116~122に演算を分割する。デバイス116~122は、演算を行い、妥当な場合には、それぞれの出力またはインディケーションをシステム100に返し、システム100は、要求された出力またはインディケーションをクライアント102に返し得る。

【0034】

ニューラルネットワーク演算を行う任意のデバイス、例えば、デバイス116~122は、命令およびデータを記憶するためのメモリ、例えば、ランダムアクセスメモリ(RAM)と、記憶されている命令を実行するためのプロセッサとを含み得る。一般的に、各デバイスは、他のデバイスから独立した演算を行うハードウェアリソースである。例えば、各デバイスは、それ自身の処理装置を有し得る。デバイスは、グラフィック処理装置(GPU)または中央処理装置(CPU)であり得る。例として、1つの機械が、1つまたは複数のデバイス、例えば、複数のCPUおよびGPUをホストし得る。

【0035】

各デバイスはまた、それぞれの計算性能を有し得る。すなわち、デバイスは、異なる量のメモリ、処理速度、または他のアーキテクチャ上の特性を有し得る。そのため、いくつかのデバイスは、他のデバイスが行うことができない演算を行い得る。例えば、いくつかの演算は、特定のデバイスのみが有するある量のメモリを要求する、またはいくつかのデバイスは、特定のタイプの演算、例えば、推論演算のみを行うように構成される。

【0036】

10

20

30

40

50

システム100内のセッションマネージャ104は、計算グラフの演算を行う間のセッションを開始する要求をクライアント102から受信する。セッションマネージャ104は、計算グラフの演算を行い得る、デバイスのセット、例えば、デバイス116～122を管理し、演算を行うのに利用可能デバイスのセットに配分器108を提供し得る。

【0037】

配分器108は、計算グラフ内で行われることになる各演算について、演算を行う、それぞれのターゲットデバイス、例えば、デバイス116を決定し、いくつかの実施形態においては、それぞれのターゲットデバイスが演算を行うための時間を決定する。他の演算が計算グラフ内の以前の演算が完了するのを要求する、例えば、他の演算が入力として以前の演算の出力を処理する一方で、いくつかの演算は並列に行われ得る。

10

【0038】

デバイスが配分器108によって割り振られた演算を行って出力を生成した後に、実行器106は、出力を読み出し得る。実行器106は、要求に対する適切な応答、例えば、出力または処理を完了し終えたということについてのインディケーションを生成し得る。その後、実行器106は、応答をクライアント102に返し得る。

【0039】

セッションマネージャ104はまた、計算グラフにおいて行われることになる演算のセットを実行器106に提供する。実行器106は、演算のグラフ実行に関連するデバイス116～122から実行時の統計値を定期的に取り出す。実行器106は、実行時の統計値を配分器108に提供し、配分器108は、さらなる演算についての配分およびスケジューリングを再最適化し得る。この再最適化については図2を参照して以下でさらに説明している。

20

【0040】

図2は、計算グラフを処理するための例示的なプロセス200のフロー図である。便宜上、プロセス200は、1つまたは複数のロケーションに位置する1つまたは複数のコンピュータのシステムによって行われるものとして説明している。例えば、計算グラフシステム、例えば、適切にプログラムされた、図1の計算グラフシステム100が、プロセス200を行い得る。

【0041】

システムが、計算グラフを処理する要求をクライアントから受信する(ステップ202)。例えば、要求は、図1を参照して上述したように、計算グラフによって表されるニューラルネットワーク推論を指定された入力に対して行う要求、計算グラフによって表されるニューラルネットワーク訓練演算を訓練データの指定されたセットに対して行う要求、または計算グラフによって表される他のニューラルネットワーク演算を行う要求であり得る。

30

【0042】

システムが、計算グラフを表すデータを取得する(ステップ204)。いくつかのケースにおいては、データは、クライアントからの要求とともに送信される。他のケースにおいては、要求は、計算グラフを特定し、システムは、特定されたグラフを表すデータをメモリから読み出す。例として、グラフを表すデータは、グラフ内のノードの配列であり得る。各ノードは、演算タイプ、名称、およびノードに対する入力方向および出力方向エッジのリストを指定する情報を含み得る。

40

【0043】

システムが要求された演算を行うための複数の利用可能なデバイスを特定する(ステップ206)。システムは、例えば、データセンタにおいて、多くのデバイスと接続し得る。システムは、例えば、図1の実行器106を使用して、各デバイスの状態を維持し得る。各デバイスは、ビジーまたは利用可能ないずれかであり得る。デバイスが他の演算を現時点行っておりさらなる演算を割り当てることができない場合には、またさなければ、グラフ処理演算を行うのに利用不可である場合には、デバイスはビジーである。デバイスにさらなる演算を割り当てることができる場合には、例えば、さらなる演算をデバイスによってキューに入れることができる場合には、デバイスは利用可能である。

【0044】

50

システムが、複数のサブグラフに計算グラフを分割する(ステップ208)。各サブグラフは、計算グラフ内の1つまたは複数のノードを含む。いくつかの実施形態においては、クライアントからの要求は、どのように計算グラフを所定のサブグラフに分割すべきかを指定するラベルを含む。例えば、ユーザは、計算グラフのためのラベルを手動で生成し、要求にラベルを含め得る。要求がそのようなラベルを含む場合には、システムは、計算グラフを所定のサブグラフに分割する。

【0045】

いくつかの他の実施形態においては、システムは、どのように計算グラフを配置するかに基づいて計算グラフを分割する。具体的には、システムは、チェーン構造内に配置されている計算グラフ内の1つまたは複数のノードを接続する有向エッジを特定するためにグラフを分析し得る。チェーン構造内のノードは、ノードからノードへの1つの有向エッジに従って互いに接続されるノードである。そのため、チェーン内のノードは、それ自身の演算を計算する前に、チェーン内の前段のノードにおける演算が計算を終了するのを待機する必要がある。サブグラフの分割については図3を参照してさらに説明している。

【0046】

さらに他の実施形態においては、システムは、グラフ内のノードをクラスタリングし、その後、同一のクラスタ内のノードを同一のサブグラフに割り当てる。具体的には、システムは、有向エッジ上を流れる共有データに対して演算を行うノードを特定するためにグラフを分析し得る。例えば、複数のノードは、入力として、前段のノードから同一のデータを受信し得る。システムは、サブグラフが特定のデバイスに割り当てられた際にノードによって表される複数の演算のために同一のデータを記憶するメモリをデバイスが再利用することができるように、同一のサブグラフ内で同一のデータを受信する、そのようなノードをクラスタリングし得る。このことについては図3を参照してさらに説明している。

【0047】

どのようにシステムがサブグラフを生成するかについてのより詳細については以下で見つけることができる。

【0048】

システムは、各サブグラフについて、サブグラフ内の1つまたは複数のノードによって表される演算をそれぞれの利用可能なデバイスに割り当てる(ステップ210)。いくつかの実施形態においては、システムは、各サブグラフを、サブグラフ内のノードによって表される演算を行うために必要な計算性能を有するデバイスに割り当てる。いくつかの実施形態においては、クライアントからの要求は、特定のノードのための演算を行う特定のタイプのデバイスを特定する、ユーザによって指定されたデータを含む。例えば、数学的に処理の重い演算を有する特定のノードはGPUに割り当てられるべきであると、ユーザは指定することができる。システムは、特定のノードを含むサブグラフを特定のタイプを有するデバイスに割り当て得る。

【0049】

いくつかの他の実施形態においては、システムは、サブグラフ内のノードを表す演算によって消費されることになるリソースの最大量を評価することによって、デバイスにどのサブグラフを割り当てるかを決定する。例えば、システムは、サブグラフ内のいずれかのノードによって消費されることになるメモリの最大量を算出し得る。具体的には、システムは、サブグラフを横断して、サブグラフの各ノードへの各有向エッジ上および各ノードからの各有向エッジ上のテンソルの次元を算出し得る。テンソルの次元は、演算を行うのにデバイスによって消費されることになるメモリのサイズを示す。システムは、サブグラフ内を流れる最大のテンソルを記憶することが可能なメモリを有するデバイスにサブグラフを割り当て得る。

【0050】

サブグラフをデバイスに割り当てる別の実施形態については図4を参照して以下でさらに説明しており、どのようにシステムがサブグラフをデバイスに割り当てるかについてのより詳細については以下で見つけることができる。

【 0 0 5 1 】

システムが、デバイスに割り当てられたノードの演算をデバイスに行わせる(ステップ212)。いくつかの実施形態においては、システムは、演算を開始する要求を各デバイスに送信する。デバイスは、要求を受信し、それに応じて、デバイスに割り当てられたノードの演算を行うことを開始する。いくつかの実施形態においては、デバイスは、デバイスに割り当てられたノードの演算を非同期的に行う。例えば、デバイスは、キュー、非ブロッキングカーネル、またはその両方を使用して非同期的に演算を行い得る。非同期的に演算を行うことを以下で説明している。

【 0 0 5 2 】

図3は、例示的な計算グラフを図示している。例として、計算グラフシステム、例えば、図1に記載のシステム100は、入力の設定が与えられたときに計算グラフを使用して推論を計算する要求をクライアントから受信し得る。具体的には、クライアントは、ノード316の出力を要求し得る。入力の設定は、ノード302に有向エッジ上で提供され得る。

10

【 0 0 5 3 】

システムは、計算グラフを3つのサブグラフ318~322に分割し得る。サブグラフ318~322を生成するために、システムは、ノードのチェーンを特定するために計算グラフを分析し得る。例えば、システムは、ノード304、316の第1のチェーン、ノード302、306、310の第2のチェーン、およびノード308、312、314の第3のチェーンを特定し得る。ノードの他の可能なチェーンが考えられるが、システムは、サブグラフの数を最小化するチェーンを選択し得る。システムは、ノードのチェーンをそれぞれのサブグラフにグループ化し得る。

20

【 0 0 5 4 】

いくつかの実施形態においては、ノード306の出力が同一である場合には、システムは、ノード306、308、および310を1つのサブグラフにグループ化する。これは、ノード310および308の両方がノード306から同一の出力を受信しているためである。この場合には、ノード310および308によって表される演算を、メモリ消費を最小化するために同一のデバイス上で行う。すなわち、デバイスは、ノード310および308の両方のための演算を行う際に、ノード306からの出力を記憶する同一のメモリロケーションにアクセスし得る。

【 0 0 5 5 】

システムは、3つのサブグラフ318~322を3つのそれぞれの利用可能なデバイスに割り当て得る。第1のサブグラフ322が初期ノード302を含み、そのノードのいずれもが他のサブグラフの出力に依存していないため、システムは、第1のサブグラフ322を割り当てることによって開始してもよい。第1のサブグラフ322を割り当てると、システムは、第2のサブグラフ318を割り当て得る。第2のサブグラフ318内のノード304は、第1のサブグラフ322に割り当てられたデバイスによって算出されることになるノード302の出力を要求する。

30

【 0 0 5 6 】

いくつかの実施形態においては、システムは、ノード302によって表される演算が完了したということについてのインディケーションを受信するまで、第2のサブグラフ318を割り当ててのを待機する。このことは、現在の情報、例えば、メモリまたはデバイスの利用可能性に基づいて、サブグラフをシステムが動的に割り当ててることを可能にしており、効率を改善している。インディケーションを受信すると、システムは、ノード302の出力のサイズをハンドリングすることが可能なデバイスに第2のサブグラフ318を割り当て得る。いくつかの他の実施形態においては、システムは、ノード302および304から有向エッジ上を流れるテンソルの次元を決定するためにグラフを分析する。システムは、その後、テンソルの次元に基づいて第2のサブグラフ318を割り当て得る。すなわち、システムは、第2のサブグラフ318に対するテンソルのメモリ要件をハンドリングし得るデバイスに第2のサブグラフ318に割り当てる。

40

【 0 0 5 7 】

同様に、第3のサブグラフ320の初期ノード308は、ノード306の出力を要求する。システムは、第1のサブグラフが割り当てられたデバイスがノード306によって表される演算を完

50

了するまで、第3のサブグラフ320を割り当ててのを待機し得る。ノード306によって表される演算が完了すると、システムは、第3のサブグラフ320をそれぞれの利用可能なデバイスに割り当ててのためにノード306の出力を分析し得る。

【0058】

デバイスは、まだ完了していない入力が必要とするノードにおいては、演算を一時中止し得る、例えば、アイドル状態に遷移し得る。例えば、ノード308のための演算を行った後に、第3のサブグラフ320に割り当てられたデバイスは、ノード312のための演算を行い得る。第3のサブグラフ320に割り当てられたデバイスは、その後、ノード310からの入力を受信されたかどうかを決定する。デバイスは、デバイスが入力をノード310から受信するまで、ノード312のための演算を行うのを待機し得る。

10

【0059】

最終ノードすなわちノード316が演算を行った後に、ノードが割り当てられているデバイスは、ノードの出力またはグラフの処理が完了したということについてのインディケーションをシステムに返し得る。システムは、その後、必要に応じて、出力をクライアントに返し得る。

【0060】

図4は、例示的なプロセス400のフロー図である。便宜上、プロセス400は、1つまたは複数のロケーションに位置する1つまたは複数のコンピュータのシステムによって行われるものとして説明している。例えば、計算グラフシステム、例えば、適切にプログラムされた、図1の計算グラフシステム100が、プロセス400を行い得る。

20

【0061】

システムが、デバイスに対するサブグラフの初期割り当てを決定する(ステップ402)。システムは、貪欲法を使用してデバイスに対する初期割り当てを決定し得る。すなわち、システムは、サブグラフ内の1つまたは複数の初期ノードを分析することによって、デバイスにどのサブグラフを割り当ててかを決定する。初期ノードは、データがサブグラフ内で流れることを開始するノードである。

【0062】

いくつかの実施形態においては、システムは、初期ノードによって表される演算によって、または、初期ノードに接続されているノードによって表される演算によって、消費されることになるメモリの量を決定する。図2および図3を参照して上述したように、システムは、消費されることになるメモリの量を図2を参照して上述したように決定するために、初期ノードへのまたは初期ノードからのテンソルの次元を分析し得る。

30

【0063】

決定した量に基づいて、システムは、少なくとも決定した量のメモリを有するデバイスにサブグラフを割り当てて。後続のノードではなく初期ノードを考慮することによって、システムは、サブグラフをデバイスに迅速に割り当て得るが、例えば、割り当てられたデバイスが十分なメモリを有しておらず、そのため、サブグラフ内で表されている後続の演算を行うためにページングを実施しなければならない場合には、後続のノードは割り当てられたデバイスが効率的に処理することが可能ではなくなりかねないリソースを要求し得るため、割り当てが最適とはならない可能性がある。

40

【0064】

システムが、統計値を決定するためにデバイスによってグラフの処理をモニタする(ステップ404)。例えば、システムは、デバイスの各々についての動作時間、アイドル時間、またはその両方をモニタし得る。動作時間は、デバイスがクライアントからの要求を完了するのにかかる時間である。すなわち、システムは、各デバイスがサブグラフの割り当てられた演算を完了するのにどれくらいかかるかを測定する。システムはまた、デバイスに割り当てられたサブグラフの処理中に各デバイスアイドルが後続の演算をどれくらい待機するかを測定し得る。

【0065】

システムが、統計値を使用して初期割り当てを調整する(ステップ406)。具体的には、

50

システムは、動作時間もしくはアイドル時間、またはその両方を最小化するように初期割り当てを調整し得る。例として、システムは、第1および第2のサブグラフのそれぞれの初期ノードに基づいて、第1のサブグラフのための演算を行う第1のデバイスと第2のサブグラフのための演算を行う第2のデバイスとをまず割り当て得る。演算を行う時間をトラッキングした後に、システムは、第1のデバイスと第2のデバイスとの間のリソースの利用を比較し得る。第1のデバイスが第2のデバイスより長い期間アイドルである場合には、今のところ、第1のデバイスは、第2のデバイスより多くの処理能力およびメモリを有しており、システムは、第1および第2のサブグラフを使用する演算のための後続の要求について第2のデバイスに対する第1のサブグラフの割り当てと第1のデバイスに対する第2のサブグラフの割り当てとを調整し得る。

10

【0066】

システムが、調整した割り当てに従ってデバイスにサブグラフを再割り当てする(ステップ408)。すなわち、上記の説明に続いて、第1および第2のサブグラフする演算のための後続の要求に応じて、システムは、第2のデバイスに第1のサブグラフを割り当てて、第1のデバイスに第2のサブグラフを割り当てる。

【0067】

システムは、割り当てを連続的に更新してパフォーマンスを改善するためにステップ404~408を繰り返し得る。例えば、システムは、アイドル時間を最小化する割り当ての調整について複数の考えられる候補が存在すると決定する場合がある。システムは、特定のサブグラフを多数の異なるデバイスに割り当てるオプションを有し得る。ある特定のサブグラフの後続の演算において、システムは、第1の考えられる候補を選択し、演算の完了に対して第1の動作時間を測定する。別の後続の演算において、システムは、第2のイテレーションの間は第2の考えられる候補を選択し、完了に対して第2の動作時間を測定する。さらに別の後続の演算において、システムは、最短の動作時間を有する考えられる候補を選択するとともに、異なるサブグラフに対する割り当てについて異なる考えられる候補を選択し得る。いくつかの実施形態においては、システムは、改善についての閾値量に達するまで前記ステップを繰り返し得る。

20

【0068】

デバイスがそれぞれのサブグラフに割り当てられた後に、デバイスは、それぞれのサブグラフの演算を行って、例えば、計算グラフによって表されるニューラルネットワーク(または他の機械学習モデル)を使用してモデル入力を処理する。演算を完了すると、デバイスは、演算が完了したことを、または、もしあれば、演算の出力を、システムに通知し得る。システムによって受信された要求は、いくつかのケースにおいては、計算グラフ内の特定のノードの1つまたは複数の出力を含む応答を指定し得る。システムは、特定のデバイスが割り当てられた1つまたは複数のデバイスから、演算が完了した後に特定のノードの出力を受信し得る。システムは、その後、図1を参照して上述したように、出力をクライアントに提供し得る。

30

【0069】

いくつかの実施形態においては、ユーザは、計算グラフの一部、例えば、計算グラフのサブグラフ、計算グラフ内のノード、または計算グラフ内の複数のノードの異なるコレクションを、他の計算グラフのコンポーネントとして再利用することができる関数として、指定することができる。具体的には、これらの実施形態においては、計算グラフを特定するシステムデータを提供した後に、ユーザは、再利用可能な関数として計算グラフの特定の一部を指定して再利用可能な関数を、関数名、例えば、システム生成された識別子またはユーザ指定の論理的な名称と関連付ける、要求を送信し得る。システムは、その後、特定の一部内のノードおよびエッジを特定するとともにその一部を関数名と関連付けるデータを保存し得る。その後に、システムは、関数に対するリファレンス、例えば、他の計算グラフ内の特定のノードの出力が関数名を有する関数に対する入力として提供されるべきであるとともに関数の出力が他の計算グラフ内の別の特定のノードに対する入力として提供されるべきであるということについてのインディケーションを含む別の計算グラフを処

40

50

理する要求を受信し得る。要求に応じて、システムは、関数名に関連付けられたグラフの一部を特定し得るし、適切な位置にそのグラフの一部を含む拡張計算グラフを生成し得る。システムは、その後、上述したように、拡張計算グラフを処理し得る。そのため、ユーザは、毎度これらの演算を表すグラフの一部を再生成する必要もなく、ある共通の再利用される演算、例えば、特定の構成のニューラルネットワーク層の演算をそれらの計算グラフ内に容易に含めることができる。

【 0 0 7 0 】

本明細書において説明した発明特定事項の実施形態および機能的な動作は、本明細書において開示した構造およびそれらの構造的均等物含む、デジタル電子回路、有形に具現化されたコンピュータソフトウェアもしくはファームウェア、コンピュータハードウェア、またはその組合せのうちの1つまたは複数で実装され得る。本明細書において説明した発明特定事項の実施形態は、1つまたは複数のコンピュータプログラム、すなわち、データ処理装置によってまたはデータ処理装置の動作を制御するために、実行のために有形非一時的プログラムキャリア上で符号化されたコンピュータプログラム命令の1つまたは複数のモジュールとして実装され得る。あるいはまたは加えて、プログラム命令は、人工的に生成された伝搬信号、例えば、データ処理装置によって実行に適した受信機装置への伝送のための情報を符号化するために生成される、機械生成された電気、光学、または電磁気信号上に、符号化され得る。コンピュータ記憶媒体は、機械可読ストレージデバイス、機械可読ストレージ回路基板、ランダムもしくはシリアルアクセスメモリデバイス、またはその組合せのうちの1つまたは複数であり得る。しかしながら、コンピュータ記憶媒体は、伝搬信号ではない。

【 0 0 7 1 】

「データ処理装置」という用語は、例として、プログラマブルプロセッサ、コンピュータ、またはマルチブルプロセッサまたはコンピュータを含む、データを処理するためのすべての種類の装置、デバイス、および機械を含む。装置は、特殊用途ロジック回路、例えば、FPGA(フィールドプログラマブルゲートアレイ)またはASIC(特定用途向け集積回路)を含み得る。装置はまた、ハードウェアに加えて、当該のコンピュータプログラムのための実行環境を作成するコード、例えば、プロセッサファームウェア、プロトコルスタック、データベース管理システム、オペレーティングシステム、またはそれらの組合せのうちの1つまたは複数を含むコードを含み得る。

【 0 0 7 2 】

コンピュータプログラム(プログラム、ソフトウェア、ソフトウェアアプリケーション、モジュール、ソフトウェアモジュール、スクリプト、またはコードとしても称され得るまたは開示され得る)は、コンパイル型もしくはインタプリタ型言語、または宣言型もしくは手続き型言語を含む、任意の形式のプログラミング言語で書かれ得るし、スタンドアロンプログラムとして、または、モジュール、コンポーネント、サブルーチン、またはコンピューティング環境における使用に適した他のユニットとしてといったことを含む、任意の形式でデプロイされ得る。コンピュータプログラムは、ファイルシステム内のファイルに対応してもよい必ずしも対応する必要はない。プログラムは、他のプログラムまたはデータ、例えば、マークアップ言語のドキュメントに記憶された1つまたは複数のスクリプトを保持するファイルの一部に、当該のプログラム専用の単一のファイルに、または、複数の協調ファイル、例えば、1つまたは複数のモジュール、サブプログラム、またはコードの一部を記憶するファイルに、記憶され得る。コンピュータプログラムは、1つのコンピュータ上で、または、1箇所に位置するもしくは複数のサイトにわたって分散され通信ネットワークによって相互通信する複数のコンピュータ上で、実行されるようにデプロイされ得る。

【 0 0 7 3 】

本明細書において使用しているように、「エンジン」または「ソフトウェアエンジン」は、入力とは異なる出力を提供する入力/出力システムを実装するソフトウェアを指す。エンジンは、機能性、例えば、ライブラリ、プラットフォーム、ソフトウェア開発キット

(「SDK」)、またはオブジェクトなどの符号化ブロックであり得る。各エンジンは、1つまたは複数のプロセッサとコンピュータ可読媒体とを含む、任意の適切なタイプのコンピュータデバイス、例えば、サーバ、モバイル電話、タブレットコンピュータ、ノートブックコンピュータ、音楽プレーヤ、電子書籍リーダー、ラップトップもしくはデスクトップコンピュータ、PDA、スマートフォン、または他の固定またはポータブルデバイスに実装され得る。加えて、2つ以上のエンジンが、同一のコンピュータデバイス上にまたは異なるコンピュータデバイス上に実装され得る。

【0074】

本明細書において説明したプロセスおよびロジックフローは、入力データを処理して出力を生成することによって機能を実施するために1つまたは複数のコンピュータプログラムを実行する1つまたは複数のプログラマブルコンピュータによって実施され得る。プロセスおよびロジックフローはまた、特殊用途ロジック回路、例えば、FPGA(フィールドプログラマブルゲートアレイ)またはASIC(特定用途向け集積回路)によって実施され得るし、装置はまた、そのような特殊用途ロジック回路として実装され得る。

【0075】

コンピュータプログラムの実行に適したコンピュータは、例として、汎用もしくは特殊用途マイクロプロセッサまたはその両方、または任意の他の種類の中央処理装置を含むまたは基づき得る。一般的に、中央処理装置は、リードオンリーメモリもしくはランダムアクセスメモリまたはその両方から命令およびデータを受信することになる。コンピュータの必須要素は、命令を実施または実行するための中央処理装置と、命令およびデータを記憶するための1つまたは複数のメモリデバイスとである。一般的に、コンピュータはまた、例えば、磁気、光磁気ディスク、または光ディスクといったデータを記憶するための1つまたは複数のマスストレージデバイスを含む、または、そのような1つまたは複数のマスストレージデバイスからデータを受信もしくはそのような1つまたは複数のマスストレージデバイスにデータを送信またはその両方を行うことが動作可能に接続される。しかしながら、コンピュータは、必ずしもそのようなデバイスを有する必要はない。さらに、コンピュータは、別のデバイス、数例挙げるとすれば、例えば、モバイル電話、携帯情報端末(PDA)、モバイルオーディオもしくはビデオプレーヤ、ゲームコンソール、グローバルポジショニングシステム(GPS)受信機、または、例えばユニバーサルシリアルバス(USB)フラッシュドライブといったポータブルストレージデバイスに組み込まれ得る。

【0076】

コンピュータプログラム命令およびデータを記憶するのに適したコンピュータ可読媒体は、例として、例えば、EPROM、EEPROM、およびフラッシュメモリデバイスといった半導体メモリデバイス、例えば、内部ハードディスクまたはリムーバブルディスクといった磁気ディスク、光磁気ディスク、ならびにCD-ROMおよびDVD-ROMディスクを含む、すべての形式の不揮発性メモリ、メディア、およびメモリデバイスを含む。プロセッサおよびメモリは、特殊用途ロジック回路によって補完され得る、または、特殊用途ロジック回路に組み込まれ得る。

【0077】

ユーザとのインタラクションを提供するために、本明細書において説明した発明特定事項の実施形態は、例えば、CRT(陰極線管)モニタ、LCD(液晶ディスプレイ)モニタ、またはOLEDディスプレイといった、ユーザに情報を表示するための表示デバイスと、例えば、キーボード、マウス、またはプレゼンス感知型ディスプレイもしくは他のサーフェスといった、コンピュータに入力を提供するための入力デバイスとを有する、コンピュータに実装され得る。他の種類のデバイスも同様に、ユーザとのインタラクションを提供するために使用され得るし、例えば、ユーザに提供されるフィードバックは、例えば、視覚フィードバック、聴覚フィードバック、または触覚フィードバックなどといった、任意の形式の感覚フィードバックであり得るし、ユーザからの入力、音響、音声、または触覚入力を含む、任意の形式で受信され得る。加えて、コンピュータは、ユーザによって使用されるデバイスにリソースを送信するとともにユーザによって使用されるデバイスからリソースを

受信することによって、例えば、ウェブブラウザから受信した要求に応じてユーザのクライアントデバイス上のウェブブラウザにウェブページを送信することによって、ユーザとのインタラクションを行い得る。

【 0 0 7 8 】

本明細書において説明した発明特定事項の実施形態は、例えば、データサーバとして、バックエンドコンポーネントを含む、または、例えば、アプリケーションサーバといった、ミドルウェアコンポーネントを含む、例えば、ユーザがそれを介して本明細書において説明した発明特定事項の実施形態とインタラクションを行い得る、グラフィックユーザインターフェースもしくはウェブブラウザを有するクライアントコンピュータといった、フロントエンドコンポーネントを含む、コンピューティングシステムにおいて実装され得る、または、そのようなバックエンド、ミドルウェア、またはフロントエンドコンポーネントのうちの1つまたは複数の任意の組合せで実装され得る。システムのコンポーネントは、任意の形式または媒体のデジタルデータ通信によって相互接続され得る、例えば、通信ネットワークによって相互接続され得る。通信ネットワークの例としては、ローカルエリアネットワーク(「LAN」)およびワイドエリアネットワーク(「WAN」)を含む、例えば、インターネットを含む。

10

【 0 0 7 9 】

コンピューティングシステムは、クライアントおよびサーバを含み得る。クライアントおよびサーバは、一般的に互いにリモートに存在しており、通信ネットワークを介して通常はインタラクションを行う。クライアントとサーバとの関係は、それぞれのコンピュータ上で動作するとともに互いにクライアント-サーバの関係を有するコンピュータプログラムによって生じる。

20

【 0 0 8 0 】

本明細書は多くの特定の実施形態詳細を含んでいるが、これらは、任意の発明の範囲または主張される可能性がある範囲を限定するものとして解釈されるべきではなく、特定の発明の特定の実施形態に特有であり得る特徴の説明として解釈されるべきである。別個の実施形態に関連して本明細書に説明したある特徴はまた、単一の実施形態における組合せで実施され得る。反対に、単一の実施形態に関連して説明した様々な特徴はまた、複数の実施形態で別々にまたは任意の適切なサブコンビネーションで実施され得る。さらに、特徴を、ある組合せで動作するように上述しているとしても、たとえそのようにはじめは主張していたとしても、いくつかのケースでは、主張した組合せのうちの1つまたは複数の特徴を、組合せから削除し得るし、主張した組合せは、サブコンビネーションまたはサブコンビネーションの変形を対象とし得る。

30

【 0 0 8 1 】

同様に、動作を特定の順序で図面に図示しているが、このことを、望ましい結果を達成するためには、図示した特定の順序でもしくはシーケンシャルな順序でそのような動作を行う必要がある、または、図示した動作をすべて行う必要がある、と理解すべきではない。ある環境においては、マルチタスク処理およびパラレル処理が有利となる場合もある。さらに、上述した実施形態における様々なシステムモジュールおよびコンポーネントの分離はすべての実施形態においてそのような分離が必要であると理解すべきではないし、説明したプログラムコンポーネントおよびシステムは一般的に単一のソフトウェア製品内に一緒に統合され得るまたは複数のソフトウェア製品にパッケージされ得ると理解すべきである。

40

【 0 0 8 2 】

上に述べたように、発明特定事項の特定の実施形態を説明してきた。他の実施形態も以下の特許請求の範囲の範囲内にある。例えば、特許請求の範囲に記載のアクションは、異なる順序で行われ、望ましい結果をそれでも達成し得る。一例として、添付の図面に図示したプロセスは、望ましい結果を達成するために、図示した特定の順序またはシーケンシャルな順序を必ずしも必要としているわけでない。ある実施形態においては、マルチタスク処理およびパラレル処理が有利となる場合もある。

50

【符号の説明】

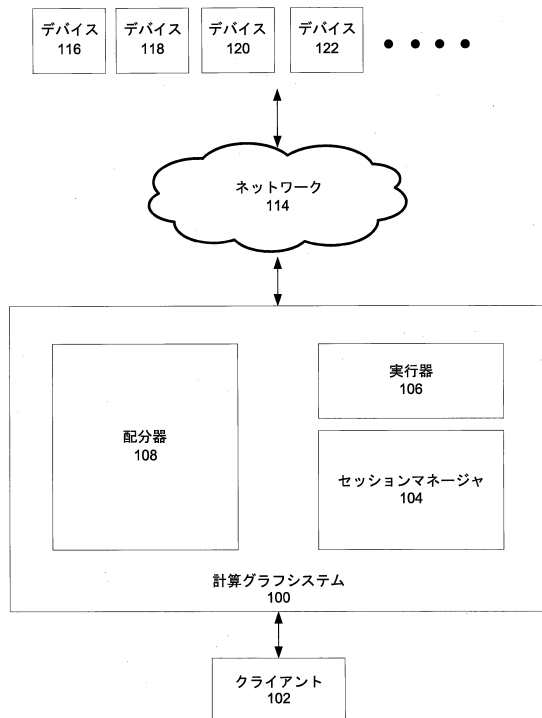
【 0 0 8 3 】

100	システム
102	クライアント
104	セッションマネージャ
106	実行器
108	配分器
114	ネットワーク
116	デバイス
118	デバイス
120	デバイス
122	デバイス
302	ノード
304	ノード
306	ノード
308	ノード
310	ノード
312	ノード
316	ノード
318	サブグラフ
320	サブグラフ
322	サブグラフ

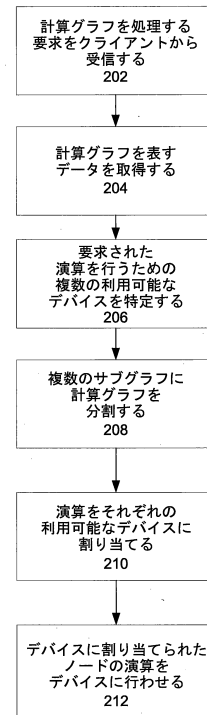
10

20

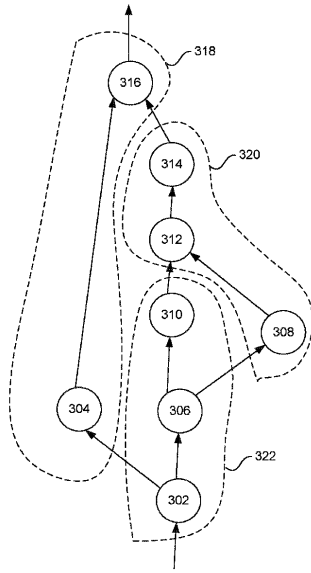
【図 1】



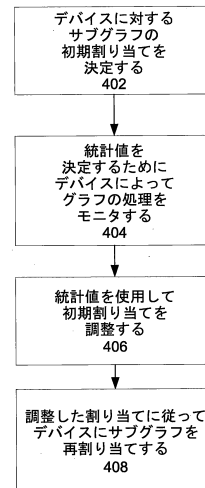
【図 2】



【図 3】



【図 4】



フロントページの続き

(31)優先権主張番号 62/253,009

(32)優先日 平成27年11月9日(2015.11.9)

(33)優先権主張国・地域又は機関
米国(US)

(72)発明者 ポール・エー・タッカー

アメリカ合衆国・カリフォルニア・94043・マウンテン・ビュー・アンフィシアター・パーク
ウェイ・1600

(72)発明者 ジェフリー・アドゲート・ディーン

アメリカ合衆国・カリフォルニア・94043・マウンテン・ビュー・アンフィシアター・パーク
ウェイ・1600

(72)発明者 サンジェイ・ゲマワット

アメリカ合衆国・カリフォルニア・94043・マウンテン・ビュー・アンフィシアター・パーク
ウェイ・1600

(72)発明者 ユアン・ユ

アメリカ合衆国・カリフォルニア・94043・マウンテン・ビュー・アンフィシアター・パーク
ウェイ・1600

合議体

審判長 田中 秀人

審判官 須田 勝巳

審判官 山澤 宏

(56)参考文献 特開平05-108595(JP,A)

特開平04-211858(JP,A)

米国特許出願公開第2010/0325621(US,A1)

特開2004-185271(JP,A)

Jeffrey Dean, et al., Large Scale Distributed
Deep Networks, Proceeding NIPS'12 of the 25
th International Conference on Neural Netw
ork Information Processing System, ACM, 2012年, V
olume 1, p.1-9, URL, <http://papers.nips.cc/paper/4687-large-scale-distributed-deep-networks.pdf>

青柳洋一、上原 稔、森 秀樹、異粒度系における静的評価に基づくタスク割付方式、情報処理
学会研究報告、日本、社団法人情報処理学会、1997年 8月22日、Vol.97, No.
78, p.7-12(97-PRO-14-2), ISSN 0919-6072

(58)調査した分野(Int.Cl., DB名)

G06N 3/10

G06F 9/50