

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号

特許第6829327号
(P6829327)

(45) 発行日 令和3年2月10日 (2021.2.10)

(24) 登録日 令和3年1月25日 (2021.1.25)

(51) Int. Cl. F I
G O 6 N 3/10 (2006.01) G O 6 N 3/10

請求項の数 13 (全 26 頁)

(21) 出願番号	特願2019-564530 (P2019-564530)	(73) 特許権者	519411401
(86) (22) 出願日	平成31年3月29日 (2019.3.29)		カンブリコン テクノロジーズ コーポレ イション リミティド
(65) 公表番号	特表2020-532778 (P2020-532778A)		中華人民共和国, ペイジン 100190
(43) 公表日	令和2年11月12日 (2020.11.12)		, ハイジャン ディストリクト, コーシュ エユアン サウス ロード, ナンバー 6
(86) 国際出願番号	PCT/CN2019/080510		, サイエントフィック リサーチ ビル ディング, スイート 644
(87) 国際公開番号	W02020/029592	(74) 代理人	100099759
(87) 国際公開日	令和2年2月13日 (2020.2.13)		弁理士 青木 篤
審査請求日	令和1年11月18日 (2019.11.18)	(74) 代理人	100123582
(31) 優先権主張番号	201810913895.6		弁理士 三橋 真二
(32) 優先日	平成30年8月10日 (2018.8.10)	(74) 代理人	100114018
(33) 優先権主張国・地域又は機関	中国 (CN)		弁理士 南山 知広

最終頁に続く

(54) 【発明の名称】 変換方法、装置、コンピューターデバイス及び記憶媒体

(57) 【特許請求の範囲】

【請求項 1】

モデル変換方法であって、

初期オフラインモデルとコンピューターデバイスのハードウェア属性情報を獲得するステップと、

前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致するかどうかを判断するステップと、及び

前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致しない場合、前記コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するステップとを含む、前記モデル変換方法。

【請求項 2】

前記コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するステップは、

前記コンピューターデバイスのハードウェア属性情報に従って、前記ターゲットオフラインモデルのモデル属性情報を決定するステップと、

前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデ

10

20

ル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つの前記モデル変換ルールをターゲットモデル変換ルールとして選択するステップと、及び

前記初期オフラインモデルのモデル属性情報と前記ターゲットモデル変換ルールに従って、前記初期オフラインモデルを前記ターゲットオフラインモデルに変換するステップとを含むことを特徴とする

請求項 1 に記載のモデル変換方法。

【請求項 3】

前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つの前記モデル変換ルールをターゲットモデル変換ルールとして選択するステップは、

10

前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つ以上の使用可能なモデル変換ルールを選択するステップと、及び

前記一つ以上の使用可能なモデル変換ルールに優先順位を付け、優先順位の高い使用可能なモデル変換ルールを前記ターゲット変換ルールとして使用するステップとを含むことを特徴とする

請求項 2 に記載のモデル変換方法。

【請求項 4】

前記一つ以上の使用可能なモデル変換ルールに優先順位を付けるステップは、

各前記使用可能なモデル変換ルールを使用して、前記初期オフラインモデルを前記ターゲットオフラインモデルに変換するためのプロセスパラメーターをそれぞれ獲得し、ここで、前記プロセスパラメーターは、変換速度、消費電力、メモリ使用率及びディスク I/O 占有率中の一つまたは複数を含むステップと、及び

20

各前記使用可能なモデル変換ルールのプロセスパラメーターに従って、一つ以上の前記使用可能なモデル変換ルールに優先順位を付けるステップとを含むことを特徴とする

請求項 3 に記載のモデル変換方法。

【請求項 5】

前記初期オフラインモデルのモデル属性情報、前記ターゲットオフラインモデルのモデル属性情報及び前記使用可能なモデル変換ルールの三つの間は 1 対 1 に対応するように保存されることを特徴とする

30

請求項 3 に記載のモデル変換方法。

【請求項 6】

前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致するかどうかを判断するステップは、

前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報に従って、前記コンピューターデバイスが前記初期オフラインモデルの実行をサポートできること決定する場合、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致すると判定するステップと、

40

前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報に従って、前記コンピューターデバイスが前記初期オフラインモデルの実行をサポートしないと決定する場合、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致しないと判定するステップとを含むことを特徴とする

請求項 1 ～ 5 のいずれか一項に記載のモデル変換方法。

【請求項 7】

前記初期オフラインモデルを獲得するステップは、

前記コンピューターデバイス上のアプリケーションソフトウェアを介して前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報を獲得するステップを含むことを特徴とする

50

請求項 1 ~ 5 のいずれか一項に記載のモデル変換方法。

【請求項 8】

前記ターゲットオフラインモデルを前記コンピューターデバイスの第 1 メモリまたは第 2 メモリに保存するステップをさらに含むことを特徴とする

請求項 1 ~ 5 のいずれか一項に記載のモデル変換方法。

【請求項 9】

前記ターゲットオフラインモデルを獲得し、前記ターゲットオフラインモデルを実行し、ここで、前記ターゲットオフラインモデルは、オリジナルネットワークにおける各オフラインノードに対応するネットワークの重み、命令及び各オフラインノードと前記オリジナルネットワーク中の他の計算ノードとの間のインターフェースデータを含むステップをさらに含むことを特徴とする

請求項 8 に記載のモデル変換方法。

【請求項 10】

モデル変換装置であって、

初期オフラインモデルとコンピューターデバイスのハードウェア属性情報を獲得するために使用される獲得モジュールと、

前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致するかどうかを判断するように使用される判断モジュールと、及び

前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致しない場合、前記コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するように使用される変換モジュールとを含む、前記モデル変換装置。

【請求項 11】

コンピュータプログラムが保存されるメモリとプロセッサを含むコンピューターデバイスであって、前記プロセッサが前記コンピュータプログラムを実行する場合、請求項 1 ~ 7 のいずれか一項に記載の方法を実現する、前記コンピューターデバイス。

【請求項 12】

前記プロセッサは、演算ユニットとコントローラーユニットを含み、前記演算ユニットは、メイン処理回路と複数のスレーブ処理回路を含み、

前記コントローラーユニットは、データ、機械学習モデル及び計算命令を獲得するために使用され、

前記コントローラーユニットは、さらに前記計算命令を解析して複数の演算命令を獲得し、前記複数の演算命令及び前記データを前記メイン処理回路に送信するように使用され、

前記メイン処理回路は、前記データ及び前記メイン処理回路と前記複数のスレーブ処理回路との間で伝送されるデータと演算命令に対して前処理を実行するように使用され、

前記複数のスレーブ処理回路は、前記メイン処理回路から伝送されたデータ及び演算命令に従って、中間演算を実行して複数の中間結果を獲得し、複数の中間結果を前記メイン処理回路に伝送するように使用され、

前記メイン処理回路は、さらに前記複数の中間結果に対して後続処理を実行して、前記計算命令の計算結果を獲得するように使用されることを特徴とする

請求項 11 に記載のコンピューターデバイス。

【請求項 13】

コンピュータプログラムが保存されたコンピュータ可読記憶媒体であって、前記コンピュータプログラムは、プロセッサによって実行される場合、請求項 1 ~ 7 のいずれか一項に記載の方法のステップを実行する、前記コンピュータ可読記憶媒体。

【発明の詳細な説明】

【技術分野】

【0001】

関連出願

本願発明は、出願日が2018年08月10日であり、出願番号が201810913895.6であり、名前が「変換方法、装置、コンピューターデバイス及び記憶媒体」である中国特許出願に対して優先権を主張し、当該中国特許出願の全ての内容は参照により本明細書に組み込まれる。

【0002】

本発明は、コンピュータ技術分野に関し、具体的に、変換方法、装置、コンピューターデバイス及び記憶媒体に関する。

10

【背景技術】

【0003】

人工知能技術の開発により、ディープラーニングはユビキタスで不可欠になり、例えば、TensorFlow、MXNet、Caffe及びPyTorchなどの多くのスケラブルなディープラーニングシステムが生成され、上記ディープラーニングシステムは、CPUやGPUなどのプロセッサで実行することができるニューラルネットワークモデルまたは他の機械学習モデルを提供するために使用されることができる。ニューラルネットワークモデルなどの機械学習モデルが異なるプロセッサで実行される必要がある場合、多くの場合、ニューラルネットワークモデルなどの機械学習モデルを変換する必要がある。

20

【0004】

従来の変換方法は次の通りである。開発者は、同じニューラルネットワークモデルに複数の異なる変換モデルを配置することができ、複数の変換モデルは、異なるプロセッサに適用されることができ、実際の使用において、コンピューターデバイスは、当該ニューラルネットワークモデル及びそれに対応する複数の変換モデルを同時に受信する必要があり、ユーザーは現在のコンピューターデバイスのタイプに従って、上記の複数のモデルから一つのモデルを選択して、当該ニューラルネットワークモデルが現在のコンピューターデバイスで実行できるようにする必要がある。しかし、上記変換方法のデータ入力量とデータ処理量はともに大きく、上記比較的により大きいデータ入力量及びデータ処理量は、コンピューターデバイスの記憶容量及び処理制限を容易に超えてしまい、コンピューターデバイスの処理速度が低下し、さらに正常に動作できないことがある。

30

【発明の概要】

【発明が解決しようとする課題】

【0005】

関連技術に存在する問題を少なくともある程度克服するために、本出願は、変換方法、装置、コンピューターデバイス及び記憶媒体を提供する。

【課題を解決するための手段】

【0006】

本出願は、初期オフラインモデルとコンピューターデバイスのハードウェア属性情報を獲得するステップと、前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致するかどうかを判断するステップと、及び前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致しない場合、前記コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するステップとを含むモデル変換方法を提供する。

40

【0007】

一つの実施例において、前記コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピューターデ

50

バイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するステップは、前記コンピューターデバイスのハードウェア属性情報に従って、前記ターゲットオフラインモデルのモデル属性情報を決定するステップと、前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つの前記モデル変換ルールをターゲットモデル変換ルールとして選択するステップと、及び前記初期オフラインモデルのモデル属性情報と前記ターゲットモデル変換ルールに従って、前記初期オフラインモデルを前記ターゲットオフラインモデルに変換するステップとを含む。

【0008】

一つの実施例において、前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つの前記モデル変換ルールをターゲットモデル変換ルールとして選択するステップは、前記初期オフラインモデルのモデル属性情報と前記ターゲットオフラインモデルのモデル属性情報に従って、複数の前記予め設定されたモデル変換ルールから一つ以上の使用可能なモデル変換ルールを選択するステップと、及び前記一つ以上の使用可能なモデル変換ルールに優先順位を付け、優先順位の高い使用可能なモデル変換ルールを前記ターゲット変換ルールとして使用するステップとを含む。

【0009】

一つの実施例において、前記一つ以上の使用可能なモデル変換ルールに優先順位を付けるステップは、各前記使用可能なモデル変換ルールを使用して、前記初期オフラインモデルを前記ターゲットオフラインモデルに変換するためのプロセスパラメーターをそれぞれ獲得し、ここで、前記プロセスパラメーターは、変換速度、消費電力、メモリ使用率及びディスクI/O占有率中の一つまたは複数を含むステップと、及び各前記使用可能なモデル変換ルールのプロセスパラメーターに従って、一つ以上の前記使用可能なモデル変換ルールに優先順位を付けるステップとを含む。

【0010】

一つの実施例において、前記初期オフラインモデルのモデル属性情報、前記ターゲットオフラインモデルのモデル属性情報及び前記使用可能なモデル変換ルールの三つの間は1対1に対応するように保存される。

【0011】

一つの実施例において、前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致するかどうかを判断するステップは、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報に従って、前記コンピューターデバイスが前記初期オフラインモデルの実行をサポートできること決定する場合、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致すると判定するステップと、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報に従って、前記コンピューターデバイスが前記初期オフラインモデルの実行をサポートしないと決定する場合、前記初期オフラインモデルのモデル属性情報と前記コンピューターデバイスのハードウェア属性情報が一致しないと判定するステップとを含む。

【0012】

一つの実施例において、前記初期オフラインモデルを獲得するステップは、前記コンピューターデバイス上のアプリケーションソフトウェアを介して前記初期オフラインモデルと前記コンピューターデバイスのハードウェア属性情報を獲得するステップを含む。

【0013】

一つの実施例において、前記ターゲットオフラインモデルを前記コンピューターデバイスの第1メモリまたは第2メモリに保存するステップをさらに含む。

【0014】

一つの実施例において、前記方法は、前記ターゲットオフラインモデルを獲得し、前記

10

20

30

40

50

ターゲットオフラインモデルを実行し、ここで、前記ターゲットオフラインモデルは、オリジナルネットワークにおける各オフラインノードに対応するネットワークの重み、命令及び各オフラインノードと前記オリジナルネットワーク中の他の計算ノードとの間のインターフェースデータを含むステップをさらに含む。

【0015】

本出願は、初期オフラインモデルとコンピュータデバイスのハードウェア属性情報を獲得するために使用される獲得モジュールと、前記初期オフラインモデルと前記コンピュータデバイスのハードウェア属性情報に従って、前記初期オフラインモデルのモデル属性情報と前記コンピュータデバイスのハードウェア属性情報が一致するかどうかを判断するように使用される判断モジュールと、及び前記初期オフラインモデルのモデル属性情報と前記コンピュータデバイスのハードウェア属性情報が一致しない場合、前記コンピュータデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、前記初期オフラインモデルを前記コンピュータデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するように使用される変換モジュールとを含むモデル変換装置を提供する。

10

【0016】

本出願は、コンピュータプログラムが保存されるメモリとプロセッサを含み、前記プロセッサが前記コンピュータプログラムを実行する場合、上記の変換方法を実現するコンピュータデバイスを提供する。

【0017】

20

一つの実施例において、前記プロセッサは、演算ユニットとコントローラユニットを含み、前記演算ユニットは、メイン処理回路と複数のスレーブ処理回路を含み、前記コントローラユニットは、データ、機械学習モデル及び計算命令を獲得するために使用され、前記コントローラユニットは、さらに前記計算命令を解析して複数の演算命令を獲得し、前記複数の演算命令及び前記データを前記メイン処理回路に送信するように使用され、前記メイン処理回路は、前記データ及び前記メイン処理回路と前記複数のスレーブ処理回路との間で伝送されるデータと演算命令に対して前処理を実行するように使用され、前記複数のスレーブ処理回路は、前記メイン処理回路から伝送されたデータ及び演算命令に従って、中間演算を実行して複数の中間結果を獲得し、複数の中間結果を前記メイン処理回路に伝送するように使用され、前記メイン処理回路は、さらに前記複数の中間結果に対して後続処理を実行して、前記計算命令の計算結果を獲得するように使用される。

30

【0018】

本出願は、コンピュータプログラムが保存され、前記コンピュータプログラムは、プロセッサによって実行される場合、上記の変換方法のステップを実行するコンピュータ可読記憶媒体を提供する。

【0019】

以上の一般的な説明及び以下の詳細な説明は単なる例示及び解釈であり、本発明を限定する意図ではないことを理解されたい。

【図面の簡単な説明】

【0020】

40

【図1】一つの実施例におけるコンピュータデバイスの構造ブロック図である。

【図2】図1におけるプロセッサの一つの実施例の構造ブロック図である。

【図3】図1におけるプロセッサの一つの実施例の構造ブロック図である。

【図4】図1におけるプロセッサの一つの実施例の構造ブロック図である。

【図5】一つの実施例におけるモデル変換方法のフロー模式図である。

【図6】一つの実施例におけるモデル変換方法のフロー模式図である。

【図7】一つの実施例におけるオフラインモデル生成方法のフロー模式図である。

【図8】一つの実施例におけるオフラインモデル生成方法のフロー模式図である。

【図9】一実施例のネットワークモデルのネットワーク構造図である。

【図10】図9におけるネットワークモデルのオフラインモデル生成プロセス模式図であ

50

る。

【発明を実施するための形態】

【0021】

本明細書の図面は、本明細書に組み込まれ、その一部を構成し、本出願の実施例を例示し、明細書とともに本出願の原理を説明するために使用される。

【0022】

ここで、例示的な実施例を詳細に説明し、その例を添付の図面に示す。以下の説明が図面に関する場合、特に明記しない限り、異なる図面の同じ数字は、同じまたは類似の要素を示す。以下の例示的な実施例で説明される実施形態は、本出願と一致するすべての実施形態を表すものではない。代わりに、それらは、添付の特許請求の範囲に詳述されている、本出願のいくつかの態様と一致する装置及び方法の単なる例である。

10

【0023】

図1は、一実施例におけるコンピューターデバイスのブロック図であり、当該コンピューターデバイスは、携帯電話またはタブレットコンピュータなどのモバイル端末、またはデスクトップコンピュータ、タブレットカード、またはクラウドサーバなどの端末であり得る。当該コンピューターデバイスは、ロボット、プリンター、スキャナー、運転レコーダー、ナビゲーター、カメラ、ビデオカメラ、プロジェクター、時計、モバイルストレージ、ウェアラブルデバイス、乗り物、家電製品、および/または医療機器に適用できる。ここで、乗り物は、飛行機、船、および/または車両を含むことができ、家電は、テレビ、エアコン、電子レンジ、冷蔵庫、炊飯器、加湿器、洗濯機、電灯、ガスコンロ、レンジフードを含むことができ、医療機器は、核磁気共鳴装置、B超音波および/または心電計などを含むことができる。

20

【0024】

当該コンピューターデバイスは、プロセッサ100、当該プロセッサ100に接続される第1メモリ200及び第2メモリ300を含むことができる。選択的に、プロセッサ100は、CPU(Central Processing Unit、中央プロセッサ)、GPU(Graphics Processing Unit、グラフィックプロセッサ)またはDSP(Digital Signal Processing、デジタル信号処理)のような汎用プロセッサであることができ、当該プロセッサ100は、IPU(Intelligence Processing Unit、インテリジェントプロセッサ)などの専用ニューラルネットワークプロセッサであることもできる。もちろん、当該プロセッサは、命令セットプロセッサ、関連チップセット、専用マイクロプロセッサ(例えば、専用集積回路(ASIC))またはキャッシュ(cache)用途に使用されるオンボードメモリなどであることもである。

30

【0025】

選択的に、図2に示したように、当該プロセッサ100は、コントローラーユニット110と演算ユニット120を含むことができ、ここで、コントローラーユニット110は、演算ユニット120に接続され、当該演算ユニット120は、一つのメイン処理回路121と複数のスレーブ処理回路122を含むことができる。コントローラーユニット110は、データ、機械学習モデル及び計算命令を獲得するために使用される。当該機械学習モデルは、具体的にネットワークモデルを含むことができ、当該ネットワークモデルは、ニューラルネットワークモデル及び/または非ニューラルネットワークモデルであることができる。コントローラーユニット110は、さらに獲得された計算命令を解析して演算命令を獲得し、複数の演算命令及びデータをメインプロセッサ回路に送信するために使用される。メイン処理回路は、データ及び当該メイン処理回路と複数のスレーブ処理回路との間で伝送されるデータ及び演算命令に対して前処理を実行するために使用される。複数のスレーブ処理回路は、メイン処理回路から伝送されたデータ及び演算命令に従って、中間演算を並列に実行して複数の中間結果を獲得し、複数の中間結果をメイン処理回路に伝送するために使用され、メイン処理回路は、さらに複数の中間結果に対して後続処理を実行して計算命令の計算結果を獲得するために使用される。

40

50

【 0 0 2 6 】

選択的に、当該コントローラユニット 1 1 0 は、命令記憶ユニット 1 1 1、命令処理ユニット 1 1 2 及び記憶キューユニット 1 1 4 を含むことができ、命令記憶ユニット 1 1 1 は、機械学習モデルに関連する計算命令を保存するために使用され、命令処理ユニット 1 1 2 は、計算命令に対して解析して複数の演算命令を獲得するために使用され、記憶キューユニット 1 1 4 は、命令キューを保存するために使用され、当該命令キューは、当該キューの前後順序によって実行される複数の演算命令または計算命令を含む。選択的に、当該コントローラユニット 1 1 0 は、複数の演算命令を備える場合、第 1 演算命令と第 1 演算命令との間の第 0 演算命令に関連関係があるかどうかを決定するための依存関係処理ユニット 1 1 3 をさらに含み、例えば、第 1 演算命令と第 0 演算命令に関連関係が存在する場合、第 1 演算命令を命令記憶ユニット内にキャッシュし、第 0 演算命令が実行完了した後、命令記憶ユニットから第 1 演算命令を抽出して演算ユニットに伝送する。具体的に、依存関係処理ユニット 1 1 3 が第 1 演算命令に従って、第 1 演算命令における必要なデータ（例えば、マトリックス）の第 1 ストレージアドレス区間を抽出し、第 0 演算命令に従って、第 0 演算命令における必要なマトリックスの第 0 ストレージアドレス区間を抽出し、例えば、第 1 ストレージアドレス区間と第 0 ストレージアドレス区間に重なる領域がある場合、第 1 演算命令と第 0 演算命令に関連関係があると決定し、例えば、第 1 ストレージアドレス区間と第 0 ストレージアドレス区間に重なる領域がない場合、第 1 演算命令と第 0 演算命令に関連関係がないと決定する。

10

【 0 0 2 7 】

一つの実施例において、図 3 に示したように、演算ユニット 1 2 0 は、分岐処理回路 1 2 3 をさらに含むことができ、ここで、メイン処理回路 1 2 1 は分岐処理回路 1 2 3 に接続され、分岐処理回路 1 2 3 は複数のスレーブ処理回路 1 2 2 に接続され、分岐処理回路 1 2 3 は、メイン処理回路 1 2 1 とスレーブ処理回路 1 2 2 との間のデータまたは命令を転送するために使用される。この実施例において、メイン処理回路 1 2 1 は、具体的に一つの入力ニューロンを複数のデータブロックに割り当て、複数のデータブロックにおける少なくとも一つの少なくとも一つのデータブロック、重み及び複数の演算命令における少なくとも一つの演算命令を分岐処理回路に送信するために使用され、分岐処理回路 1 2 3 は、メイン処理回路 1 2 1 と複数のスレーブ処理回路 1 2 2 との間のデータブロック、重み及び演算命令を転送するために使用され、複数のスレーブ処理回路 1 2 2 は、当該演算命令に従って受信されたデータブロック及び重みに対して演算して中間結果を獲得し、中間結果を分岐処理回路 1 2 3 に伝送するために使用され、メイン処理回路 1 2 1 は、さらに分岐処理回路によって送信された中間結果を後続処理して当該計算命令の結果を獲得し、当該計算命令の結果を前記コントローラユニットに送信するために使用される。

20

30

【 0 0 2 8 】

他の選択可能な実施例において、図 4 に示したように、演算ユニット 1 2 0 は、一つのメイン処理回路 1 2 1 と複数のスレーブ処理回路 1 2 2 を含むことができる。ここで、複数のスレーブ処理回路は、アレイ配置され、それぞれのスレーブ処理回路は隣接する他のスレーブ処理回路に接続され、メイン処理回路は、複数のスレーブ処理回路における k 個のスレーブ処理回路に接続され、k 個のスレーブ処理回路は、第 1 行の n 個のスレーブ処理回路、第 m 行の n 個のスレーブ処理回路及び第 1 列の m 個のスレーブ処理回路であり、説明すべきのは、図 1 C に示した K 個のスレーブ処理回路は、第 1 行の n 個のスレーブ処理回路、第 m 行の n 個のスレーブ処理回路及び第 1 列の m 個のスレーブ処理回路のみを含み、即ち、当該 k 個のスレーブ処理回路は、複数のスレーブ処理回路におけるメイン処理回路に直接に接続されるスレーブ処理回路である。K 個のスレーブ処理回路は、メイン処理回路と複数のスレーブ処理回路との間でデータ及び命令を転送するために使用される。

40

【 0 0 2 9 】

選択的に、上記のメイン処理回路 1 2 1 は、変換処理回路、起動処理回路、加算処理回路中の一つまたは任意の組合せを含むことができ、変換処理回路は、メイン処理回路が受信したデータブロックまたは中間結果によって第 1 データ構造と第 2 データ構造との間の

50

交換（例えば、連続データと離散データの変換）を実行するために使用され、またはメイン処理回路が受信したデータブロックまたは中間結果によって第１データタイプと第２データタイプとの間の交換（例えば、固定小数点タイプと浮動小数点タイプの変換）を実行するために使用され、起動処理回路は、メイン処理回路内のデータの起動演算を実行するために使用され、加算処理回路は、加算演算または累積演算を実行するために使用される。

【００３０】

当該第１メモリ２００または第２メモリ３００は本出願の実施例で提供されるモデル変換方法を実行するためのコンピュータプログラムを保存することができる。具体的に、当該モデル変換方法は、コンピューターデバイスが受信した初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致しない場合、当該コンピューターデバイスが当該初期オフラインモデルを実行することができるように、当該初期オフラインモデルを当該コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換する。上記のモデル変換方法において、コンピューターデバイスのデータ入力量が少なく、人間の介入なしで初期オフラインモデルからターゲットオフラインモデルへの変換を自動的に完了することができ、変換プロセスが簡単で、変換効率が低い。さらに、第１メモリ２００または第２メモリ３００に保存されたコンピュータプログラムは、さらに本出願の実施例で提供されるオフラインモデル生成方法を実現するために使用される。

【００３１】

選択的に、第１メモリ２００は、ネットワーク入力データ、ネットワーク出力データ、ネットワークの重み及び命令などのネットワークモデルの実行プロセスにおける関連データを保存するために使用されることができる。当該第１メモリ２００は、キャッシュなどの揮発性メモリのような内部メモリであることができる。当該第２メモリ３００は、ネットワークモデルに対応するオフラインモデルを保存するために使用されることができ、例えば、第２メモリ３００は、不揮発性メモリであることができる。

【００３２】

上記コンピューターデバイスの動作原理は、以下のモデル変換方法における各ステップの実行プロセスと一致し、当該コンピューターデバイスのプロセッサ１００がメモリ２００におけるコンピュータプログラムを実行すると、モデル変換方法における各ステップを実現し、具体的には、以下の説明を参照することができる。

【００３３】

もちろん、他の実施例において、当該コンピューターデバイスは、プロセッサと一つのメモリをさらに含み、当該コンピューターデバイスは、プロセッサが当該プロセッサに接続されたメモリを含むことができる。当該プロセッサは、図２～４に示したプロセッサを使用することができ、その具体的な構造は、上記のプロセッサ１００に関する説明を参照することができる。当該メモリは、複数の記憶ユニットを含むことができ、例えば、当該メモリは、第１記憶ユニット、第２記憶ユニット及び第３記憶ユニットを含むことができ、ここで、当該第１記憶ユニットは、コンピュータプログラムを保存するために使用されることができ、当該コンピュータプログラムは、本出願の実施例で提供されるモデル変換方法を実現するために使用される。当該第２記憶ユニットは、オリジナルネットワーク実行プロセスにおける関連データを保存するために使用されることができ、当該第３記憶ユニットは、オリジナルネットワークに対応するオフラインモデル及び当該オフラインモデルに対応するターゲットオフラインモデル及び予め設定されたモデル変換ルールなどを保存するために使用されることができる。さらに、当該メモリに含まれる記憶ユニットの数は、３より多いこともでき、ここでは具体的に限定しない。

【００３４】

図５に示したように、本出願の実施例は、上記のコンピューターデバイスに適用可能なモデル変換方法を提供する。当該モデル変換方法は、コンピューターデバイスが受信した初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属

性情報が一致しない場合、当該コンピューターデバイスが当該初期オフラインモデルを実行することができるように、当該初期オフラインモデルを当該コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換する。具体的に、上記方法は、以下のステップを含むことができる。

【 0 0 3 5 】

S 1 0 0 において、初期オフラインモデル及びコンピューターデバイスのハードウェア属性情報を獲得する。

【 0 0 3 6 】

具体的に、上記の初期オフラインモデルとは、コンピューターデバイスが直接に獲得したオフラインモデルを指し、当該初期オフラインモデルは、第 2 メモリ 3 0 0 に保存されることができる。ここで、オフラインモデルは、オリジナルネットワークにおける計算ノードのネットワークの重み及び命令などの必要なネットワーク構造情報を含むことができ、命令は、当該計算ノードがどのような計算機能を実行するかを示すために使用され、具体的に当該オリジナルネットワークにおける各計算ノードの計算属性及び各計算ノードの間の接続関係などの情報を含むことができる。当該オリジナルネットワークは、ニューラルネットワークモデルなどのネットワークモデルであることができ、図 9 を参照することができる。従って、コンピューターデバイスは、当該オリジナルネットワークに対応するオフラインモデルを実行することによって、当該オリジナルネットワークの演算機能を実現することができ、同じオリジナルネットワークを再びコンパイルする必要がないため、当該ネットワークを実行するプロセッサの実行時間を短縮し、さらにプロセッサの処理速度及び効率を向上させる。選択的に、本出願の実施例における初期オフラインモデルは、オリジナルネットワークに従って直接に生成されたオフラインモデルであることができ、他のモデル属性情報を有するオフラインモデルの一回または複数回の変換後に獲得されたオフラインモデルであることもでき、ここでは具体的に限定しない。

【 0 0 3 7 】

S 2 0 0 において、初期オフラインモデルとコンピューターデバイスのハードウェア属性情報に従って、初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致するかどうかを判断する。

【 0 0 3 8 】

具体的に、当該初期オフラインモデルのモデル属性情報は、当該初期オフラインモデルのデータ構造及びデータタイプなどを含むことができる。例えば、初期オフラインモデルにおける各ネットワークの重み及び命令の配置方法、各ネットワークの重みのタイプ及び各命令のタイプなどである。上記のコンピューターデバイスのハードウェア属性情報は、当該コンピューターデバイスのモデル番号、当該コンピューターデバイスが処理できるデータタイプ（例えば、固定小数点と浮動小数点など）及びデータ構造などを含む。コンピューターデバイスは、上記初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報に従って、当該コンピューターデバイスが上記初期オフラインモデルの実行をサポートできるかどうかを判断して（即ち、当該コンピューターデバイスが上記の初期オフラインモデルを実行できるかどうか）、当該初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致するかどうかを決定する。

【 0 0 3 9 】

選択的に、当該コンピューターデバイスが上記初期オフラインモデルの実行をサポートできる場合、当該初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致すると判定する。当該コンピューターデバイスが上記初期オフラインモデルの実行をサポートしない場合、当該初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致しないと判定する。

【 0 0 4 0 】

例えば、当該コンピューターデバイスは、そのハードウェア属性情報に従って、当該コンピューターデバイスが処理できるデータタイプ及びデータ構造を決定することができ、

当該初期オフラインモデルのモデル属性情報に従って、当該初期オフラインモデルのデータタイプ及びデータ構造を決定する。当該コンピューターデバイスが自身のハードウェア属性情報及び上記初期オフラインモデルのモデル属性情報に従って、当該コンピューターデバイスが上記初期オフラインモデルのデータタイプ及びデータ構造をサポートできると決定する場合、当該初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致すると判定することができる。当該コンピューターデバイスが自身のハードウェア属性情報及び上記の初期オフラインモデルのモデル属性情報に従って、当該コンピューターデバイスが上記初期オフラインモデルのデータタイプ及びデータ構造をサポートしないと決定する場合、当該初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致しないと判定することができる。

10

【 0 0 4 1 】

初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致しない場合、ステップ S 3 0 0 を実行する：コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、初期オフラインモデルをコンピューターデバイスのハードウェア属性情報に一致するターゲットオフラインモデルに変換する。

【 0 0 4 2 】

具体的に、当該初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致しない場合、当該コンピューターデバイスが当該初期オフラインモデルを使用して対応する演算を実現するように、当該コンピューターデバイスが当該初期オフラインモデルの実行をサポートできないことを示し、このとき、当該初期オフラインモデルに変換処理を行う必要がある。即ち、当該初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致しない場合、コンピューターデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、上記の初期オフラインモデルを当該コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換することができる。ここで、当該ターゲットオフラインモデルのモデル属性情報は当該コンピューターデバイスのハードウェア属性情報と一致し、当該コンピューターデバイスは、当該ターゲットオフラインモデルの実行をサポートすることができる。上記のように、当該ターゲットオフラインモデルは、オリジナルネットワークにおける各計算ノードのネットワークの重み及び命令などの必要なネットワーク構造情報も含む。

20

30

【 0 0 4 3 】

選択的に、上記の予め設定されたモデル変換ルールは、コンピューターデバイスの第 1 メモリまたは第 2 メモリに事前に保存されることができる。初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致しない場合、初期オフラインモデルをターゲットオフラインモデルに変換するために、コンピューターデバイスは、第 1 メモリまたは第 2 メモリから対応するモデル変換ルールを獲得することができる。選択的に、上記予め設定されたモデル変換ルールは、一つ以上であることができ、一つ以上のモデル変換ルールは、上記の初期オフラインモデル及びターゲットオフラインモデルと 1 対 1 に対応するように保存されることができる。例えば、上記の初期オフラインモデル、ターゲットオフラインモデル及び予め設定されたモデル変換ルールは、マッピングテーブルの方法によって、対応して保存されることができる。選択的に、初期オフラインモデルのモデル属性情報、ターゲットオフラインモデルのモデル属性情報及び使用可能なモデル変換ルールの三つの間は 1 対 1 に対応するように保存される。

40

【 0 0 4 4 】

初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致する場合、ステップ S 4 0 0 を実行することができる：コンピューターデバイスは、受信した初期オフラインモデルを直接に実行することができ、即ち、コンピューターデバイスは、オリジナルネットワークの演算機能を実現するために、当該初期オフライ

50

ンモデルに含まれるネットワークの重み及び命令に従って演算を実行することができる。具体的に、当該初期オフラインモデルのモデル属性情報と当該コンピューターデバイスのハードウェア属性情報が一致する場合、当該コンピューターデバイスが当該初期オフラインモデルの実行をサポートできることを示し、このとき、当該初期オフラインモデルに変換処理を行う必要がなく、コンピューターデバイスは、当該初期オフラインモデルを実行することができる。本実施例において、当該初期オフラインモデルを直接に実行するとは、当該初期オフラインモデルを使用して当該オリジナルネットワークに対応する機械ラーニングアルゴリズム（例えば、ニューラルネットワークアルゴリズム）を実行し、フォワード演算を実行することによりアルゴリズムのターゲットアプリケーション（例えば、音声認識などの人工知能アプリケーション）を実現することを指す。

10

【 0 0 4 5 】

本出願の実施例におけるモデル変換方法では、コンピューターデバイスは一つの初期オフラインモデルを受信するだけでよく、予め設定されたモデル変換ルールと当該コンピューターデバイスのハードウェア属性情報に従って、当該初期オフラインモデルをターゲットオフラインモデルに変換することができ、複数の異なるモデルデータを獲得する必要なく、コンピューターデバイスのデータ入力量を大幅に削減し、過剰なデータ入力量に起因するコンピューターデバイスのストレージ容量を超えるなどの問題を回避し、コンピューターデバイスの正常実行を保証する。同時に、上記モデル変換方法は、当該コンピューターデバイスのデータ処理量を削減することができ、さらに処理効率を向上させることができ、消費電力を削減する。同時に、上記のモデル変換プロセスにおいて、人間の介入が必要なく、自動化の程度が高く、使用が便利である。

20

【 0 0 4 6 】

選択的に、図 6 に示したように、上記のステップ S 3 0 0 は、以下のステップを含むことができる。

【 0 0 4 7 】

S 3 1 0 において、コンピューターデバイスのハードウェア属性情報に従って、ターゲットオフラインモデルのモデル属性情報を決定する。

【 0 0 4 8 】

具体的に、当該コンピューターデバイスは、自身のハードウェア属性情報に従って、当該コンピューターデバイスがサポートできるデータタイプ及びデータ構造などを決定することができるため、ターゲットオフラインモデルのデータタイプ及びデータ構造などのモデル属性情報を決定することができる。即ち、当該コンピューターデバイスは、そのハードウェア属性情報に従ってターゲットオフラインモデルのモデル属性情報を決定することができ、当該ターゲットオフラインモデルのモデル属性情報は、ターゲットオフラインモデルのデータ構造及びデータタイプなどの情報を含むことができる。

30

【 0 0 4 9 】

S 3 2 0 において、初期オフラインモデルのモデル属性情報とターゲットオフラインモデルのモデル属性情報に従って、複数の予め設定されたモデル変換ルールから一つのモデル変換ルールをターゲットモデル変換ルールとして選択する。

【 0 0 5 0 】

具体的に、初期オフラインモデル、ターゲットオフラインモデルと予め設定されたモデル変換ルールに 1 対 1 に対応するマッピング関係が存在するため、当該初期オフラインモデルのモデル属性情報とターゲットオフラインモデルのモデル属性情報に従って、ターゲットモデル変換ルールを決定することができる。ここで、当該モデル変換ルールは、データタイプの変換方法及びデータ構造の変換方法などを含むことができる。

40

【 0 0 5 1 】

S 3 3 0 において、初期オフラインモデルのモデル属性情報とターゲットモデル変換ルールに従って、初期オフラインモデルをターゲットオフラインモデルに変換する。

【 0 0 5 2 】

具体的に、コンピューターデバイスは、当該ターゲットモデル変換ルールによって提供

50

される変換方法に従って、初期オフラインモデルをターゲットオフラインモデルに変換して、当該コンピューターデバイスが当該ターゲットオフラインモデルを実行して演算を実行できるようにする。

【0053】

選択的に、初期オフラインモデルとターゲットオフラインモデルとの間に複数の使用可能なモデル変換ルールが存在する場合、コンピューターデバイスは、予め設定されたアルゴリズムに従ってターゲットオフラインモデルを自動的に選択することができる。具体的に、上記のステップS320は、以下のステップをさらに含む。

【0054】

初期オフラインモデルのモデル属性情報とターゲットオフラインモデルのモデル属性情報に従って、複数の予め設定されたモデル変換ルールから一つ以上の使用可能なモデル変換ルールを選択する。

10

【0055】

一つ以上の使用可能なモデル変換ルールに優先順位を付け、優先順位の高い使用可能なモデル変換ルールを前記ターゲット変換ルールとして使用する。

【0056】

具体的に、上記の使用可能なモデル変換ルールとは、初期オフラインモデルをターゲットオフラインモデルに変換する変換ルールを指す。一つ以上の使用可能なモデル変換ルールが存在する場合、初期オフラインモデルをターゲットオフラインモデルに変換する方法には、多数の異なる方法があることを示す。上記の優先順位を付ける方法は、予め設定されることができ、ユーザーによって定義されることもできる。

20

【0057】

選択的に、コンピューターデバイスは、各使用可能なモデル変換ルールに対応するプロセスパラメーターを獲得することができ、各使用可能なモデル変換ルールに対応するプロセスパラメーターに従って、一つ以上の使用可能なモデル変換ルールに優先順位を付ける。ここで、当該プロセスパラメーターは、当該使用可能なモデル変換ルールを使用して、初期オフラインモデルをターゲットオフラインモデルに変換する場合に関与するコンピューターデバイスのパフォーマンスパラメータであることができる。選択的に、当該プロセスパラメーターは、変換速度、消費電力、メモリ使用率及びディスクI/O占有率中の一つまたは複数を含むことができる。

30

【0058】

例えば、コンピューターデバイスは、初期オフラインモデルをターゲットオフラインモデルに変換するプロセスにおける変換速度、変換消費電力、メモリ使用率及びディスクI/O占有率などの一つまたは複数プロセスパラメーターの組合せに従って、上記の各使用可能なモデル変換ルールに対してスコアリングし（例えば、スコア値を取得するために、各参照要因に加重計算を行う）、最高スコアの使用可能なモデル変換ルールを最高優先順位の使用可能なモデル変換ルールとして使用することができる。即ち、最高スコアの使用可能なモデル変換ルールを当該ターゲット変換ルールとして使用することができる。その後、上記ステップS330を実行することができ、初期オフラインモデルのモデル属性情報とターゲットモデル変換ルールに従って、初期オフラインモデルをターゲットオフラインモデルに変換する。このようにして、コンピューターデバイスの処理速度及び効率を向上させるために、モデル変換プロセス中のデバイスのパフォーマンス要因を合わせて、より優れたモデル変換ルールを選択することができる。

40

【0059】

選択的に、図3に示したように、上記方法は、以下のステップをさらに含む。

【0060】

S500において、ターゲットオフラインモデルをコンピューターデバイスの第1メモリまたは第2メモリに保存する。

【0061】

具体的に、コンピューターデバイスは、獲得されたターゲットオフラインモデルを当該

50

コンピューターデバイスのローカルメモリ（例えば、第１メモリ）に保存することができる。選択的に、コンピューターデバイスは、さらに獲得されたターゲットオフラインモデルを当該コンピューターデバイスに接続された外部メモリ（例えば、第２メモリ）に保存することができ、当該外部メモリは、クラウドストレージまたは他のメモリなどであることができる。このようにして、当該コンピューターデバイスが当該ターゲットオフラインモデルを繰り返して使用する必要がある場合、上記の変換操作を繰り返して実行する必要なく、当該ターゲットオフラインモデルを直接に読み取ることで、対応する演算を実現することができる。

【００６２】

選択的に、コンピューターデバイスが上記のターゲットオフラインモデルを実行するプロセスは、以下のステップを含むことができる。

【００６３】

オリジナルネットワークを実行するために、ターゲットオフラインモデルを獲得し、ターゲットオフラインモデルを実行する。ここで、ターゲットオフラインモデルは、オリジナルネットワークにおける各オフラインノードに対応するネットワークの重み、命令及び各オフラインノードとオリジナルネットワーク中の他の計算ノードとの間のインターフェースデータなどの必要なネットワーク構造情報を含む。

【００６４】

具体的に、コンピューターデバイスが当該ターゲットオフラインモデルを使用して演算を実現する必要がある場合、第１メモリまたは第２メモリから当該ターゲットオフラインモデルを直接に獲得することができ、ターゲットオフラインモデル中のネットワークの重み及び命令などに従って演算を行って、オリジナルネットワークの演算機能を実現する。このようにして、当該コンピューターデバイスが当該ターゲットオフラインモデルを繰り返して使用する必要がある場合、上記の変換操作を繰り返す必要なく、当該ターゲットオフラインモデルを直接に読み取るだけで、対応する演算を実現することができる。

【００６５】

本実施例において、当該ターゲットオフラインモデルを直接に実行するとは、当該初期ターゲットオフラインモデルを使用して当該オリジナルネットワークに対応する機械学習アルゴリズム（例えば、ニューラルネットワークアルゴリズム）を実行し、フォワード演算を実行することによりアルゴリズムのターゲットアプリケーション（例えば、音声認識などの人工知能アプリケーション）を実現することを指す。

【００６６】

選択的に、当該コンピューターデバイスには、アプリケーションソフトウェア（Application）がインストールされ、コンピューターデバイス上のアプリケーションソフトウェアを介して、初期オフラインモデルを獲得することができ、当該アプリケーションソフトウェア（Application）は、メモリ２００または外部メモリから初期オフラインモデルを読み取るためのインターフェイスを提供することができるため、当該アプリケーションソフトウェアを介して当該初期オフラインモデルを獲得することができる。選択的に、当該アプリケーションソフトウェアは、コンピューターデバイスのハードウェア属性情報を読み取るためのインターフェイスをさらに提供するため、当該アプリケーションソフトウェアを介して当該コンピューターデバイスのハードウェア属性情報を獲得することができる。選択的に、当該アプリケーションソフトウェアは、予め設定されたモデルルールを読み取るためのインターフェイスをさらに提供することができるため、当該アプリケーションソフトウェアを介して予め設定されたモデル変換ルールを獲得することができる。

【００６７】

もちろん、他の実施例において、当該コンピューターデバイスは、入力／出力インターフェイス（例えば、Ｉ／Ｏインターフェイス）などをさらに提供ことができ、当該入力／出力インターフェイスは、初期オフラインモデルの獲得、またはターゲットオフラインモデルの出力のために使用される。当該コンピューターデバイスのハードウェア属性情

10

20

30

40

50

報は、当該コンピューターデバイスに事前に保存されることができる。

【0068】

一つの実施例において、図7に示したように、本出願の実施例は、オリジナルネットワークモデルに従って、オフラインモデルを生成する方法をさらに提供し、獲得されたオリジナルネットワークの関連データに従って、当該オリジナルネットワークのオフラインモデルを生成及び保存するために使用されるため、プロセッサが当該オリジナルネットワークを再び実行する場合、同じオリジナルネットワークを再びコンパイルする必要がなく、当該オリジナルネットワークに対応するオフラインモデルを直接に実行することができるため、当該ネットワークを実行するプロセッサの実行時間を短縮し、さらにプロセッサの処理速度及び効率を向上させる。ここで、オリジナルネットワークは、図9に示したようなネットワークのようなニューラルネットワークまたは非ニューラルネットワークなどのネットワークモデルであることができる。具体的に、上記方法は、以下のステップを含む。

10

【0069】

S010において、オリジナルネットワークのモデルデータセット及びモデル構造パラメータを獲得する。

【0070】

具体的に、コンピューターデバイスのプロセッサは、オリジナルネットワークのモデルデータセット及びモデル構造パラメータを獲得することができ、当該オリジナルネットワークのモデルデータセット及びモデル構造パラメータを介して、当該オリジナルネットワークのネットワーク構造図を獲得することができる。ここで、モデルデータセットは、オリジナルネットワークにおける各計算ノードに対応するネットワークの重みなどのデータを含み、図9に示したニューラルネットワークにおけるW1~W6は、計算ノードのネットワークの重みを表すために使用される。モデル構造パラメータは、オリジナルネットワークにおける複数の各計算ノードの接続関係及び各計算ノードの計算属性を含み、ここで、計算ノードの間の接続関係は、計算ノードの間にデータ伝達があるかどうかを表すために使用され、例えば、複数の各計算ノードの間にデータフローの伝達がある場合、複数の計算ノードの間に接続関係があると説明することができる。さらに、計算ノードの接続関係は、入力関係及び出力関係などを含むことができる。図9に示したように、計算ノードF1が計算ノードF4及びF5の入力として出力する場合、計算ノードF1と計算ノードF4との間に接続関係があり、計算ノードF1と計算ノードF4との間に接続関係があると説明できる。別の例として、計算ノードF1と計算ノードF2との間にデータ伝達がない場合、計算ノードF1と計算ノードF2との間に接続関係が存在しないと説明できる。

20

30

【0071】

各計算ノードの計算属性は、対応する計算ノードの計算タイプ及び計算パラメータを含むことができ、ここで、計算ノードの計算タイプとは、当該計算ノードがある計算を完了することに使用されることを指し、例えば、各計算ノードの計算タイプは、加算、減算、及び畳み込み算など含むことができ、対応的に、当該計算ノードは、加算を実現するための計算ノード、減算を実現するための各計算ノードまたは畳み込み算を実現するための計算ノードなどであることができる。計算ノードの計算パラメータは、当該計算ノードに対応する計算タイプを完了する必要なパラメータであることができる。例えば、計算ノードの計算タイプは、加算を実現するための計算ノードであることができ、対応的に、当該計算ノードの計算パラメータは、加算における加数であることができ、当該加算における被加数は、入力データとして獲得モジュールを介して獲得することができ、または、当該加算における被加数は、当該計算ノードの前の各計算ノードの出力データなどであることができる。

40

【0072】

S020において、オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従ってオリジナルネットワークを実行して、オリジナルネットワークにおける各計算ノードに対応する命令を獲得する。

50

【 0 0 7 3 】

具体的に、コンピュータデバイスのプロセッサは、オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従って当該オリジナルネットワークを実行して、オリジナルネットワークにおける各計算ノードに対応する命令を獲得することができる。さらに、プロセッサは、当該オリジナルネットワークの入力データをさらに獲得することができ、オリジナルネットワークの入力データ、ネットワークモデルデータセット及びモデル構造パラメータに従ってオリジナルネットワークを実行して、当該オリジナルネットワークにおける各計算ノードに対応する命令を獲得することができる。さらに、上記当該オリジナルネットワークを実行して各計算ノードの命令を獲得する過程は、実際にコンパイルの過程であり、当該コンパイル過程は、コンピュータシステムのプロセッサまたは仮想デバイスを介して実現することができる。即ち、コンピュータシステムのプロセッサまたは仮想デバイスは、オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従ってオリジナルネットワークを実行する。ここで、仮想デバイスとは、メモリのメモリ空間でプロセッサ実行空間のセクションを仮想することを指す。

10

【 0 0 7 4 】

明らかに、本実施例におけるオリジナルネットワークの実行とは、プロセッサが人工ニューラルネットワークモデルデータを使用して、ある機械学習アルゴリズム（例えば、ニューラルネットワークアルゴリズム）を実行し、フォワード演算を実行することにより、アルゴリズムのターゲットアプリケーション（例えば、音声認識などの人工知能アプリケーション）を実現することを指す。

20

【 0 0 7 5 】

S 0 3 0 において、オリジナルネットワークの各計算ノードに対応するネットワークの重み及び命令に従って、オリジナルネットワークに対応するオフラインモデルを生成し、オリジナルネットワークに対応するオフラインモデルを不揮発性メモリ（データベース）に保存する。

【 0 0 7 6 】

具体的に、当該プロセッサの制御モジュールは、オリジナルネットワークの各計算ノードに対応するネットワークの重み及び命令に従って、当該オリジナルネットワークに対応するオフラインモデルを生成することができ、例えば、当該プロセッサの制御モジュールは、オリジナルネットワークの各計算ノードに対応するネットワークの重み及び命令を不揮発性第2メモリに保存して、オフラインモデルの生成及び保存を実現することができる。ここで、オリジナルネットワークの各計算ノードについて、当該計算ノードのネットワークの重みと命令は、1対1に対応するように保存される。このようにして、当該オリジナルネットワークを再び実行する場合、当該オリジナルネットワークの各計算ノードに対してオンラインでコンパイルして命令を獲得することなく、不揮発性メモリから当該オリジナルネットワークに対応するオフラインモデルを直接に獲得することができ、それに対応するオフラインモデルに従ってオリジナルネットワークを実行し、システムの実行速度及び効率を向上させる。

30

【 0 0 7 7 】

明らかに、本実施例において、当該オリジナルネットワークに対応するオフラインモデルを直接に実行するとは、オフラインモデルを使用して当該オリジナルネットワークに対応する機械学習アルゴリズム（例えば、ニューラルネットワークアルゴリズム）実行し、フォワード演算を実行することによりアルゴリズムのターゲットアプリケーション（例えば、音声認識などの人工知能アプリケーション）を実現することを指す。

40

【 0 0 7 8 】

選択的に、図8に示したように、前記ステップS 2 0 0 は、以下のステップを含むことができる。

【 0 0 7 9 】

S 0 2 1 において、オリジナルネットワークのモデル構造パラメータに従って、オリジナルネットワークにおける各計算ノードの実行順序を獲得する。

50

【 0 0 8 0 】

具体的に、プロセッサは、オリジナルネットワークのモデル構造パラメータに従って、オリジナルネットワークにおける各計算ノードの実行順序を獲得することができ、さらに、プロセッサの演算モジュールは、オリジナルネットワークにおける各計算ノードの接続関係に従って、オリジナルネットワークにおける各計算ノードの実行順序を獲得することができる。例えば、図 9 に示したように、計算ノード F 4 の入力データは、計算ノード F 1 の出力データ及び計算ノード F 2 の出力データであり、計算ノード F 6 の入力データは、計算ノード F 4 の出力データ及び計算ノード F 5 の出力データである。従って、図 9 に示したニューラルネットワークにおける各計算ノードの実行順序は、F 1 - F 2 - F 3 - F 4 - F 5 - F 6 または F 1 - F 3 - F 2 - F 5 - F 4 - F 6 などであることができる。もちろん、計算ノード F 1、F 2 及び F 3 は、並列に実行することができ、計算ノード F 4 及び F 5 も並列に実行することができ、ここでは単に例として説明し、実行順序は具体的に限定されない。

10

【 0 0 8 1 】

S 0 2 2 において、オリジナルネットワークにおける各計算ノードの実行順序に従ってオリジナルネットワークを実行して、オリジナルネットワークにおける各計算ノードに対応する命令をそれぞれ獲得する。

【 0 0 8 2 】

具体的に、プロセッサは、オリジナルネットワークにおける各計算ノードの実行順序に従って当該オリジナルネットワークを実行して、オリジナルネットワークにおける各計算ノードに対応する命令を獲得することができ、即ち、プロセッサは、オリジナルネットワークのモデルデータセットなどのデータをコンパイルして各計算ノードに対応する命令を獲得することができ、各計算ノードに対応する命令を介して当該各計算ノードがどのようなコンピューティング機能を実現するかを知ることができ、即ち、当該各計算ノードの計算タイプ及び計算パラメータなどの計算属性を獲得することができる。

20

【 0 0 8 3 】

さらに、図 8 に示したように、前記ステップ S 3 0 0 はさらに以下のステップを含む。

【 0 0 8 4 】

S 0 3 1 において、オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従って、オリジナルネットワークのメモリ割り当て方法を獲得する。

30

【 0 0 8 5 】

具体的に、プロセッサは、オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従って、オリジナルネットワークのメモリ割り当て方法を獲得することができ、オリジナルネットワークのモデル構造パラメータに従って、オリジナルネットワークにおける各計算ノードの実行順序を獲得することができ、オリジナルネットワークにおける各計算ノードの実行順序に従って、現在のネットワークのメモリ割り当て方法を決定する。例えば、各計算ノードの実行順序に従って、各計算ノードの実行過程中的関連データは、一つのスタックに保存される。ここで、メモリ割り当て方法とは、オリジナルネットワークにおける各計算ノードに関連するデータ（入力データ、出力データ、ネットワークの重みデータ及び中間結果データなどを含む）がメモリ空間（例えば、第 1 メモリ）での保存位置を決定することを指す。例えば、データテーブルを使用して各計算ノードに関連するデータ（入力データ、出力データ、ネットワークの重みデータ及び中間結果データなど）とメモリ空間のマッピング関係を保存することができる。

40

【 0 0 8 6 】

S 0 3 2 において、オリジナルネットワークのメモリ割り当て方法に従って、オリジナルネットワークの実行過程中的関連データを第 1 メモリに保存する。ここで、オリジナルネットワークの実行過程中的関連データは、オリジナルネットワークの各計算ノードに対応するネットワークの重み、命令、入力データ、中間計算結果及び出力データなどを含む。例えば、図 9 に示したように、X 1 と X 2 は、当該ニューラルネットワークの入力データを表し、Y は、当該ニューラルネットワークの出力データを表し、プロセッサは、当該

50

ニューラルネットワークの出力データをロボットまたは異なるデジタルインターフェースを制御する制御命令に転換することができる。W 1 ~ W 6 は、計算ノード F 1、F 2 及び F 3 に対応するネットワークの重みを表すために使用され、計算ノード F 1 ~ F 5 の出力データは、中間計算結果として使用することができる。プロセッサは、決定されたメモリ割り当て方法に従って、オリジナルネットワークの実行過程中的関連データを内部メモリまたはキャッシュなどの揮発性メモリのような第 1 メモリに保存することができ、具体的な保存方法は、図 1 0 における左半部の保存空間を参照することができる。

【 0 0 8 7 】

S 0 3 3 において、第 1 メモリからオリジナルネットワークにおける各計算ノードに対応するネットワークの重み及び命令を獲得し、オリジナルネットワークにおける各計算ノードに対応するネットワークの重み及び命令を第 2 メモリに保存し、オフラインモデルを生成する。ここで、第 2 メモリは、外部メモリなどの不揮発性メモリであることができる。当該オフラインモデルの生成過程は、具体的に図 1 0 を参照することができ、図 1 0 における右半部の保存空間に保存されたのは、オリジナルネットワークの対応するオフラインモデルである。

【 0 0 8 8 】

図 9 及び図 1 0 に示したように、以下、図面に合わせて、上記オフラインモデルの生成過程を説明する。

【 0 0 8 9 】

まず、図 9 に示したように、プロセッサは、当該オリジナルネットワークのモデルデータセット、モデル構造パラメータ及び入力データを獲得することができるため、当該オリジナルネットワークのモデルデータセット及びモデル構造パラメータに従って当該オリジナルネットワークのネットワーク構造図を獲得することができる。

【 0 0 9 0 】

次に、プロセッサは、オリジナルネットワークのモデル構造パラメータに従って、オリジナルネットワークの各計算ノードの接続関係を獲得することができ、各計算ノードの接続関係に従ってオリジナルネットワークにおける各計算ノードの実行順序、及びオリジナルネットワークの実行過程中的メモリ割り当て方法を獲得するため、オリジナルネットワークの実行過程中的関連データの保存位置を獲得することができる。図 1 0 の左半部の保存空間に示したように、オリジナルネットワークの実行過程中的関連データは、各計算ノード実行順序に従って一つのスタックに保存されることができる。

【 0 0 9 1 】

最後に、プロセッサは、オリジナルネットワークの各計算ノードに対応するネットワークの重み及び命令を不揮発性の第 2 メモリに保存し、オフラインモデルを生成することができ、当該オフラインモデルの保存方法は、図 8 における右半部の保存空間を参照することができる。また、当該オフラインモデルは、当該オリジナルネットワークを実行するに必要なネットワークの重み及び命令などのデータのみを含み、オリジナルネットワークの実行過程中的の入力データ、出力データまたは中間計算結果などを保存する必要がないため、第 2 メモリにおける保存空間の消費を減少することができる。

【 0 0 9 2 】

さらなる改善として、オフラインモデルには、ノードインターフェースデータがさらに含まれ、ノードインターフェースデータは、オリジナルネットワークの各計算ノードの接続関係を表すために使用される。具体的に、ノードインターフェースデータは、各計算ノードの入力データソース及び出力データソースを含むことができる。例えば、図 9 に示したように、ノードインターフェースデータは、計算ノード F 1、F 2 及び F 3 を開始各計算ノードとして含むことができ、それぞれ予め設定された入力データを入力し、計算ノード F 1 の出力データは、計算ノード F 4 及び計算ノード F 5 の入力データなどである。このようにして、当該オリジナルネットワークを再び実行する場合、オリジナルネットワークの開始各計算ノードと入力データのみを獲得し、その後、当該オリジナルネットワークに対応するオフラインモデルに従って当該オリジナルネットワークを実行することができ

10

20

30

40

50

る。

【0093】

図5～8のフロチャートにおける各ステップは矢印で示されるように順次的に表示されるが、これらのステップは必ずしも矢印で示される順序で実行されるとは限らないことを理解されたい。本明細書で明示的に説明される場合を除き、これらのステップの実行は厳密に制限されておらず、これらのステップは他の順序で実行されてもよい。さらに、図5～8における少なくとも一部のステップは、複数のサブステップまたは複数の段階を含むことができ、これらのサブステップまたは段階は、必ずしも同時に実行される必要なく、異なる時間に実行されることができ、これらのサブステップまたは段階の実行順序も必ずしも順次に行われる必要なく、他のステップまたは他のステップのサブステップまたは段階の少なくとも一部と交代にまたは交互に行われることができる。

10

【0094】

当業者は、上記実施例方法における全部または一部のプロセスの実現は、コンピュータプログラムに介して関連するハードウェアを命令して完了し、前記コンピュータプログラムは、不揮発性コンピュータ可読記憶媒体に保存されることができ、当該プログラムは実行時に、上記各方法の実施例のプロセスを含むことができることを理解できる。ここで、本出願で提供される各実施例に使用されるメモリ、ストレージ、データベースまたは他の媒体への参照は、すべて不揮発性及び/または揮発性メモリを含むことができる。不揮発性メモリは、読み取り専用メモリ(ROM)、プログラマブルROM(PROM)、電氣的にプログラム可能なROM(EPROM)、電氣的に消去可能なプログラム可能なROM(EEPROM)またはフラッシュメモリを含むことができる。揮発性メモリは、ランダムアクセスメモリ(RAM)または外部高速キャッシュメモリを含むことができる。制限ではない説明として、RAMは、スタティックRAM(SRAM)、ダイナミックRAM(DRAM)、同期DRAM(SDRAM)、ダブルデータレートSDRAM(DDRSDRAM)、拡張SDRAM(ESDRAM)、同期リンク(Synchlink)、DRAM(SLDRAM)、ラムバス(Rambus)ダイレクトRAM(RDRAM)、ダイレクトラムバスダイナミックRAM(DRDRAM)、ラムバスダイナミックRAM(RDRAM)などのようなさまざまな形式で獲得することができる。

20

【0095】

本出願の実施例は、獲得モジュール、判断モジュール及び変換モジュールを含むモデル変換装置をさらに提供する。ここで、獲得モジュールは、初期オフラインモデルとコンピュータデバイスのハードウェア属性情報を獲得するために使用され、判断モジュールは、初期オフラインモデルとコンピュータデバイスのハードウェア属性情報に従って、初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致するかどうかを判断するように使用され、変換モジュールは、初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致しない場合、コンピュータデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、初期オフラインモデルを前記コンピュータデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換するように使用される。

30

【0096】

選択的に、当該モデル変換装置は、当該コンピュータデバイスにインストールされたアプリケーションソフトウェア(Application)であることができる。当該アプリケーションソフトウェア(Application)は、メモリ200または外部メモリから初期オフラインモデルを読み取るためのインターフェイスを提供することができるため、当該アプリケーションソフトウェアを介して当該初期オフラインモデルを獲得することができる。選択的に、当該アプリケーションソフトウェアは、コンピュータデバイスのハードウェア属性情報を読み取るためのインターフェイスをさらに提供するため、当該アプリケーションソフトウェアを介して当該コンピュータデバイスのハードウェア属性情報を獲得することができる。選択的に、当該アプリケーションソフトウェアは、予め設定されたモデルルールを読み取るためのインターフェイスをさらに提供することがで

40

50

きるため、当該アプリケーションソフトウェアを介して予め設定されたモデル変換ルールを獲得することができる。さらに、当該アプリケーションソフトウェアは、初期オフラインモデルとコンピューターデバイスのハードウェア属性情報及び予め設定されたモデル変換ルールに従って、当該初期オフラインモデルのモデル属性情報とコンピューターデバイスのハードウェア属性情報が一致しない場合、当該初期オフラインモデルを当該コンピューターデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換し、当該ターゲットオフラインモデルを第1メモリまたは第2メモリに保存する。

【0097】

モデル変換装置に関する具体的な限定は、上記のモデル変換方法に対する限定を参照することができ、ここでは繰り返して説明しない。上記モデル変換装置中の各モジュールは、ソフトウェア、ハードウェア、及びそれらの組み合わせによって全体的または部分的に実現されることができる。プロセッサが上記のモジュールに対応する操作を実行するために呼び出されるように、上記各モジュールは、コンピューター装置のプロセッサに組み込まれるか、コンピューター装置のプロセッサから独立していてもよく、ソフトウェア形式でコンピューター装置のメモリに保存されてもよい。

【0098】

なお、本発明の一実施例は、コンピュータプログラムが保存されたコンピュータ可読記憶媒体をさらに提供し、前記コンピュータプログラムは、プロセッサによって実行される場合、請求項に記載のいずれかの一つの実施例の方法のステップを実現する。選択的に、当該コンピュータ可読記憶媒体は、不揮発性及び/または揮発性メモリを含むことができる。不揮発性メモリは、読み取り専用メモリ(ROM)、プログラマブルROM(PROM)、電氣的にプログラム可能なROM(EPROM)、電氣的に消去可能なプログラム可能なROM(EEPROM)またはフラッシュメモリを含むことができる。揮発性メモリは、ランダムアクセスメモリ(RAM)または外部高速キャッシュメモリを含むことができる。制限ではない説明として、RAMは、スタティックRAM(SRAM)、ダイナミックRAM(DRAM)、同期DRAM(SDRAM)、ダブルデータレートSDRAM(DDRSDRAM)、拡張SDRAM(ESDRAM)、同期リンク(Synclink)、DRAM(SLDRAM)、ラムバス(Rambus)ダイレクトRAM(RDRAM)、ダイレクトラムバスダイナミックRAM(DRDRAM)、ラムバスダイナミックRAM(RDRAM)などのようなさまざまな形式で獲得することができる。

【0099】

具体的に、上記コンピュータプログラムがプロセッサによって実行される場合、以下のステップを実現する。

【0100】

S100において、初期オフラインモデルとコンピューターデバイスのハードウェア属性情報を獲得する。

【0101】

具体的に、上記の初期オフラインモデルとは、コンピューターデバイスが直接に獲得したオフラインモデルを指し、当該初期オフラインモデルは、第2メモリ300に保存されることができる。ここで、オフラインモデルは、オリジナルネットワークにおける計算ノードのネットワークの重み及び命令などの必要なネットワーク構造情報を含むことができ、ここで、命令は、当該計算ノードがどのような計算機能を実行するかを示すために使用され、具体的に当該オリジナルネットワークにおける各計算ノードの計算属性及び各計算ノードの間の接続関係などの情報を含むことができる。選択的に、本出願の実施例における初期オフラインモデルは、オリジナルネットワークに従って直接に生成されたオフラインモデルであることができ、他のモデル属性情報を有するオフラインモデルの一回または複数回の変換後に獲得されたオフラインモデルであることもでき、ここでは具体的に限定しない。

【0102】

S200において、初期オフラインモデルとコンピューターデバイスのハードウェア属

10

20

30

40

50

性情報に従って、初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致するかどうかを判断する。

【0103】

具体的に、当該初期オフラインモデルのモデル属性情報は、当該初期オフラインモデルのデータ構造及びデータタイプなどを含むことができる。例えば、初期オフラインモデルにおける各ネットワークの重み及び命令の配置方法、各ネットワークの重みのタイプ及び各命令のタイプなどである。上記のコンピュータデバイスのハードウェア属性情報は、当該コンピュータデバイスのモデル番号、当該コンピュータデバイスが処理できるデータタイプ（例えば、固定小数点と浮動小数点など）及びデータ構造などを含む。

【0104】

例えば、当該コンピュータデバイスは、そのハードウェア属性情報に従って、当該コンピュータデバイスが処理できるデータタイプ及びデータ構造を決定することができ、当該初期オフラインモデルのモデル属性情報に従って、当該初期オフラインモデルのデータタイプ及びデータ構造を決定する。当該コンピュータデバイスが自身のハードウェア属性情報及び上記初期オフラインモデルのモデル属性情報に従って、当該コンピュータデバイスが上記初期オフラインモデルのデータタイプ及びデータ構造をサポートできると決定する場合、当該初期オフラインモデルのモデル属性情報と当該コンピュータデバイスのハードウェア属性情報が一致すると判定することができる。当該コンピュータデバイスが自身のハードウェア属性情報及び上記の初期オフラインモデルのモデル属性情報に従って、当該コンピュータデバイスが上記初期オフラインモデルのデータタイプ及びデータ構造をサポートしないと決定する場合、当該初期オフラインモデルのモデル属性情報と当該コンピュータデバイスのハードウェア属性情報が一致しないと判定することができる。

【0105】

初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致しない場合、ステップS300を実行する：コンピュータデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、初期オフラインモデルをコンピュータデバイスのハードウェア属性情報に一致するターゲットオフラインモデルに変換する。

【0106】

具体的に、当該初期オフラインモデルのモデル属性情報と当該コンピュータデバイスのハードウェア属性情報が一致しない場合、当該コンピュータデバイスが当該初期オフラインモデルを使用して対応する演算を実現するように、当該コンピュータデバイスが当該初期オフラインモデルの実行をサポートできないことを示し、このとき、当該初期オフラインモデルに変換処理を行う必要がある。即ち、当該初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致しない場合、コンピュータデバイスのハードウェア属性情報と予め設定されたモデル変換ルールに従って、上記の初期オフラインモデルを当該コンピュータデバイスのハードウェア属性情報と一致するターゲットオフラインモデルに変換することができる。ここで、当該ターゲットオフラインモデルのモデル属性情報は当該コンピュータデバイスのハードウェア属性情報と一致し、当該コンピュータデバイスは、当該ターゲットオフラインモデルの実行をサポートすることができる。上記のように、当該ターゲットオフラインモデルは、オリジナルネットワークにおける各計算ノードのネットワークの重み及び命令などの必要なネットワーク構造情報も含む。

【0107】

初期オフラインモデルのモデル属性情報とコンピュータデバイスのハードウェア属性情報が一致する場合、ステップS400を実行することができる：コンピュータデバイスは、受信した初期オフラインモデルを直接に実行することができ、即ち、コンピュータデバイスは、オリジナルネットワークの演算機能を実現するために、当該初期オフラインモデルに含まれるネットワークの重み及び命令に従って演算を実行することができる。

具体的に、当該初期オフラインモデルのモデル属性情報と当該コンピュータデバイスのハードウェア属性情報が一致する場合、当該コンピュータデバイスが当該初期オフラインモデルの実行をサポートできることを示し、このとき、当該初期オフラインモデルに変換処理を行う必要がなく、コンピュータデバイスは、当該初期オフラインモデルを実行することができる。

【0108】

プロセッサによって上記コンピュータプログラムを実行するプロセスは、上記実施例のモデル変換方法の実行プロセスと一致することは明らかであり、具体的には、上記の説明を参照することができ、ここでは繰り返して説明しない。

【0109】

上記各実施例の同一または類似の部分は相互に参照され、いくつかの実施例で詳細に説明されていない部分は他の実施例の同一または類似の内容を参照することができることが理解される。

【0110】

本出願の説明において、「第1」、「第2」などの用語は説明目的のみに使用され、相対的な重要性を示すまたは暗示するものとして解釈されるべきではないことに留意されたい。なお、本出願の説明において、「複数」の意味は、特に明記しない限り、少なくとも二つを意味する。

【0111】

フローチャートまたはここで他の方法で説明した任意のプロセスまたは方法は、特定の論理機能またはプロセスのステップを実現するための一つまたはさらに複数の実行可能な命令を含むコードのモジュール、セグメントまたは部分を表すと理解され得る。本出願の選択的な実施例の範囲には、追加の実現が含まれ、ここで、機能は、示されたまたは議論された順序ではなく、関与される機能に応じた実質的に同時の方法または反対の順序で実行されることができ、これは、本出願の実施例が属する技術分野の当業者によって理解されるべきである。

【0112】

本出願の各部分は、ハードウェア、ソフトウェア、ファームウェア、またはそれらの組み合わせで実現されることを理解されたい。上述実施形態において、複数のステップまたは方法は、メモリに保存され、且つ適切な命令実行システムによって実行されるソフトウェアまたはファームウェアで実現されることができ。例えば、別の実施形態のようにハードウェアで実現する場合、当技術分野で周知の以下の技術のいずれか一つまたは組み合わせによって実現することができる：データ信号に論理機能を実装する論理ゲートを備えた個別論理回路、適切な組み合わせ論理ゲートを備えた専用集積回路、プログラマブルゲートアレイ(PGA)、フィールドプログラマブルゲートアレイ(FPGA)など。

【0113】

当業者は、上記実施例を実現する方法を実現するステップのすべてまたは一部は、関連するハードウェアに命令するプログラムによって完了でき、前記プログラムは、コンピュータ可読記憶媒体に保存されることができ、当該プログラムは実行される場合に、方法実施例のステップの一つまたはその組み合わせを含むことを理解できる。

【0114】

なお、本出願の各実施例における各機能ユニットは、一つの処理モジュールに統合されてもよく、各ユニットは物理的に別々に存在してもよく、または二つまたはそれ以上のユニットが一つのモジュールに統合されてもよい。上記の統合されたモジュールは、ハードウェアの形態で実現されてもよく、ソフトウェア機能モジュールの形態で実現されてもよい。前記統合されたモジュールは、ソフトウェア機能モジュールの形態で実現され、スタンドアロン製品として販売または使用される場合、コンピュータ可読記憶媒体に保存されることもできる。

【0115】

上述の記憶媒体は、読み取り専用メモリ、磁気ディスク、または光ディスクなどである

10

20

30

40

50

ことができる。

【 0 1 1 6 】

本明細書の説明において、「一実施例」、「いくつかの実施例」、「例示」、「具体的な実施例」、または「いくつかの実施例」などの説明は、当該実施例または例示に結合して説明される具体的な特徴、構造、材料または特性が、本出願の少なくとも一つの実施例または例示に含まれることを意味する。本明細書において、上記用語の概略的な表現は、必ずしも同じ実施例または例示を意味するものではない。さらに、説明される具体的な特徴、構造、材料、または特性は、任意の一つまたは複数の実施例または例示において適切な方法で結合されることができる。

【 0 1 1 7 】

以上で本出願の実施例を示して説明したが、上述の実施例は例示であり、本出願の範囲を限定するものとして解釈されるべきではなく、当業者は、本出願の範囲内で、上記実施例の変化、修正、変更、及び変形を行うことができる。

10

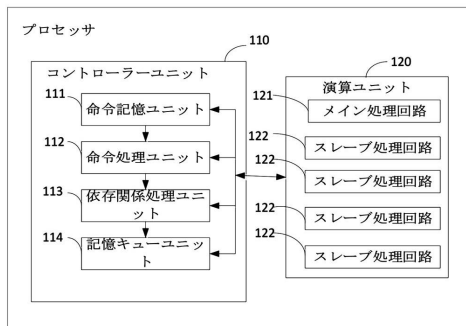
【 図 1 】

図 1



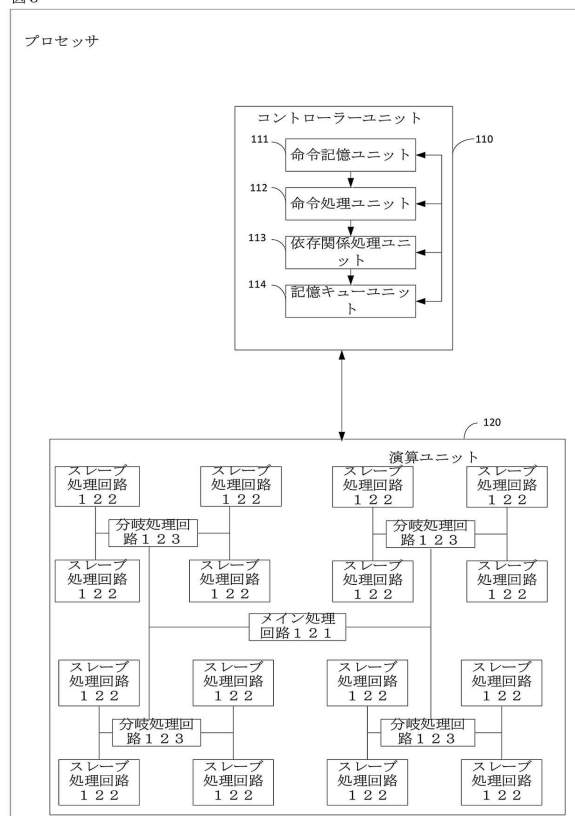
【 図 2 】

図 2

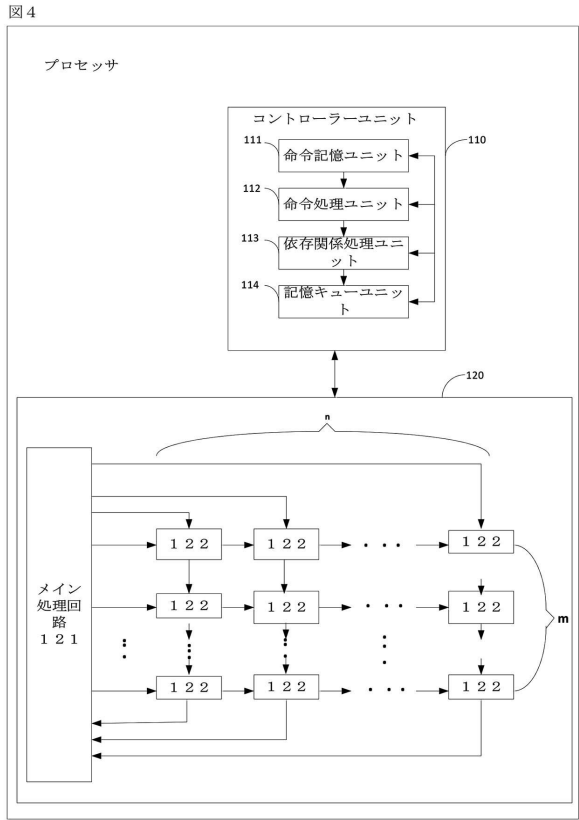


【 図 3 】

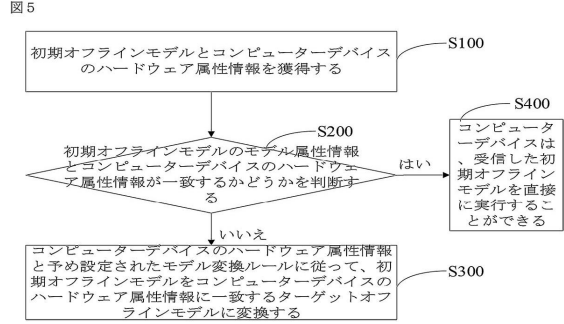
図 3



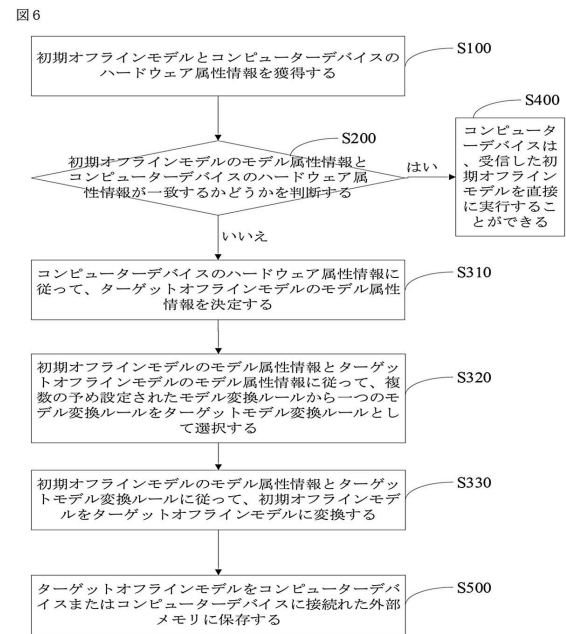
【図 4】



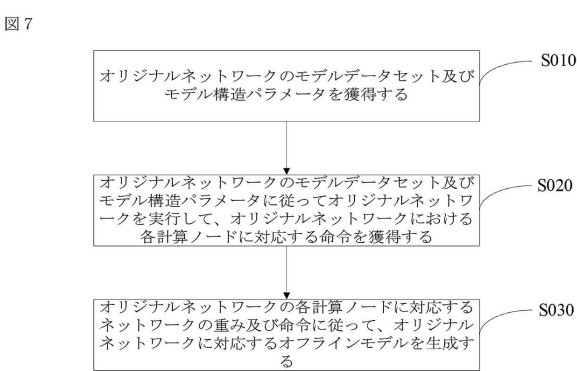
【図 5】



【図 6】

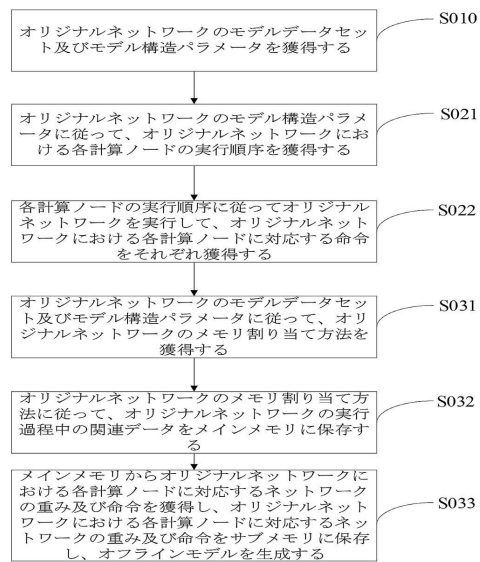


【図 7】



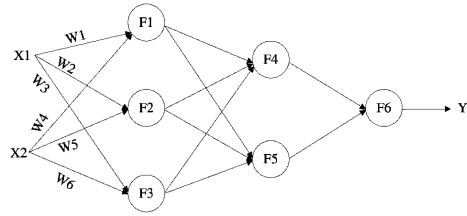
【図 8】

図 8



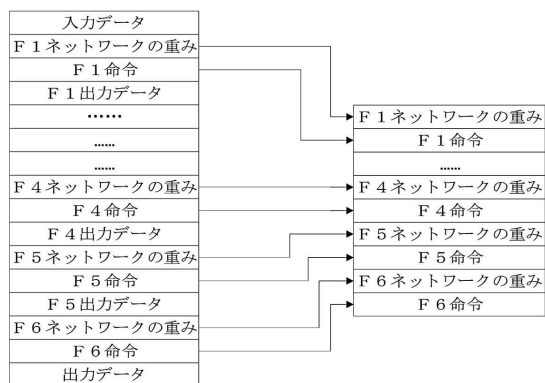
【図 9】

図 9



【図 10】

図 10



フロントページの続き

(74)代理人 100165191

弁理士 河合 章

(74)代理人 100133835

弁理士 河野 努

(72)発明者 リウ シャオリー

中華人民共和国, ペイジン 100190, ハイジャン ディストリクト, コーシュエユアン サウス ロード, ナンバー 6, サイエнтиフィック リサーチ コンプレックス ビルディング, ルーム 644

(72)発明者 リャン チュン

中華人民共和国, ペイジン 100190, ハイジャン ディストリクト, コーシュエユアン サウス ロード, ナンバー 6, サイエнтиフィック リサーチ コンプレックス ビルディング, ルーム 644

(72)発明者 クオ チー

中華人民共和国, ペイジン 100190, ハイジャン ディストリクト, コーシュエユアン サウス ロード, ナンバー 6, サイエнтиフィック リサーチ コンプレックス ビルディング, ルーム 644

審査官 久保 光宏

(56)参考文献 米国特許出願公開第2016/0328644 (US, A1)

「ソニーも参戦、深層学習ソフト 組み込み向けの開発環境で競う」, 日経エレクトロニクス, 日本, 日経BP社, 2017年 8月20日, 2017年9月号 (第1183号), 第61~64頁, ISSN: 0385-1680

@Vengineer, 「特集 知っ得 世界のAI技術 第2部第2章 グーグルTensorFlowがいろんなプロセッサに対応できるメカニズム」, インターフェース (Interface), 日本, CQ出版株式会社, 2018年 2月 1日, 2018年2月号 (第44巻, 第2号), 第64~73頁, ISSN: 0387-9569

松尾恒, 「Deep Learning推論の多階層高速化フレームワーク」, 画像ラボ, 日本, 日本工業出版株式会社, 2018年 8月10日, 2018年8月号 (第29巻, 第8号), 第29~37頁, ISSN: 0915-6755

古川新, 「ラズパイで始める人工知能 機械学習で顔認識モデルを作ろう 最終回 TensorFlowモデルの実行方法」, 日経BPパソコンベストムック ラズパイマガジン, 日本, 日経BP社, 2018年 7月 9日, 2018年8月号, 第90~94頁, ISBN: 978-4-8222-9271-3 (文献に記載の発行日: 2018年8月22日, (財)ソフトウェア情報センター受入日: 2018年7月9日)

(58)調査した分野(Int.Cl., DB名)

G06N3/00 - 99/00

CSDB (日本国特許庁)

IEEE Explore (IEEE)