

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号
特許第7613588号
(P7613588)

(45)発行日 令和7年1月15日(2025.1.15)

(24)登録日 令和7年1月6日(2025.1.6)

(51)国際特許分類

F I

G 0 6 N 20/00 (2019.01)

G 0 6 N 20/00 1 3 0

G 0 6 N 3/098(2023.01)

G 0 6 N 3/098

請求項の数 7 (全22頁)

(21)出願番号	特願2023-534452(P2023-534452)	(73)特許権者	000004237
(86)(22)出願日	令和3年7月12日(2021.7.12)		日本電気株式会社
(86)国際出願番号	PCT/JP2021/026148		東京都港区芝五丁目7番1号
(87)国際公開番号	WO2023/286129	(74)代理人	100103090
(87)国際公開日	令和5年1月19日(2023.1.19)		弁理士 岩壁 冬樹
審査請求日	令和6年1月5日(2024.1.5)	(74)代理人	100124501
			弁理士 塩川 誠人
		(72)発明者	吉山 智之
			東京都港区芝五丁目7番1号 日本電気株式会社内
		審査官	新井 則和

最終頁に続く

(54)【発明の名称】 学習システムおよび学習方法

(57)【特許請求の範囲】

【請求項1】

サーバと、複数のクライアントとを備える学習システムであって、
各クライアントは、
共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、前記重み付け和の計算に関わるパラメータとを学習する学習手段と、
前記所定の複数の操作のパラメータと前記重み付け和の計算に関わるパラメータとのうち、前記所定の複数の操作のパラメータを前記サーバに送信するクライアント側パラメータ送信手段とを備え、
前記サーバは、
前記各クライアントから受信した前記所定の複数の操作のパラメータに基づいて、前記所定の複数の操作のパラメータを計算し直すパラメータ計算手段と、
前記所定の複数の操作のパラメータを前記各クライアントに送信するサーバ側パラメータ送信手段とを備え、
前記所定の複数の操作の数は、前記複数のクライアントの数未満である
ことを特徴とする学習システム。

【請求項2】

前記各クライアントの前記学習手段は、前記重み付け和の計算に関わるパラメータを独自に学習する

請求項 1 に記載の学習システム。

【請求項 3】

前記所定の複数の操作は、いずれも線形操作である

請求項 1 または請求項 2 に記載の学習システム。

【請求項 4】

前記各クライアントは、

前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータが確定した場合に、前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータに基づいて、前記所定の複数の操作を 1 つの操作に変換する変換手段を備える請求項 3 に記載の学習システム。

10

【請求項 5】

前記各クライアントは、

前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータが確定した場合に、前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータによって定まるモデルに基づいて、与えられたデータに対する推論結果を導出する推論手段を備える

請求項 1 から請求項 4 のうちのいずれか 1 項に記載の学習システム。

【請求項 6】

サーバと、複数のクライアントとによって行われる学習方法であって、

各クライアントが、

共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、前記重み付け和の計算に関わるパラメータとを学習し、

20

前記所定の複数の操作のパラメータと前記重み付け和の計算に関わるパラメータとのうち、前記所定の複数の操作のパラメータを前記サーバに送信し、

前記サーバが、

前記各クライアントから受信した前記所定の複数の操作のパラメータに基づいて、前記所定の複数の操作のパラメータを計算し直し、

前記所定の複数の操作のパラメータを前記各クライアントに送信し、
前記所定の複数の操作の数は、前記複数のクライアントの数未満である

30

ことを特徴とする学習方法。

【請求項 7】

前記各クライアントが、前記重み付け和の計算に関わるパラメータを独自に学習する

請求項 6 に記載の学習方法。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、モデルのパラメータを学習する学習システム、学習方法、および、学習プログラム、並びに、推論器に関する。

【背景技術】

40

【0002】

一般に、機械学習では、学習データが多いほど、推論精度の高いモデルを学習することができる。そのため、複数のクライアントがそれぞれ独自にデータを保持している場合には、サーバが各クライアントからデータを集め、サーバがそのデータを学習データとして、モデルを学習することが考えられる。

【0003】

しかし、個々のクライアントの立場では、データを外部に提供することは、データ流出の観点から好ましくない。このことは、個々のクライアントが、それぞれ別々の管理者（例えば別々の会社）によって管理されている場合に、特に当てはまる。例えば、個々の会社は、独自に有しているデータを、外部に提供したくないと考える。従って、サーバが各

50

クライアントからデータを集め、サーバがそのデータを学習データとしてモデルを学習することは困難である場合が多い。

【 0 0 0 4 】

そこで、連合学習が提案されている。以下に、連合学習の一例を示す。連合学習では、例えば、サーバが得たモデル（グローバルモデルと称する。）を、各クライアントに提供する。各クライアントは、グローバルモデルと、クライアントが独自に持っているデータとに基づいて、モデルを学習する。クライアントが学習によって得るモデルをローカルモデルと記す。各クライアントは、ローカルモデル、または、グローバルモデルとローカルモデルとの差分情報を、サーバに送信する。サーバは、各クライアントから得た各ローカルモデル（または、各差分情報）に基づいて、グローバルモデルを更新し、再度、グローバルモデルを各クライアントに提供する。本例の連合学習では、上記の処理を繰り返す。例えば、サーバが、グローバルモデルを各クライアントに提供してから、サーバが、グローバルモデルを更新するまでの動作を繰り返す。そして、例えば、上記の動作の繰り返し回数が所定回数に達したことが、学習終了条件であると定めておき、上記の動作の繰り返し回数が所定回数に達したときに、サーバが得たグローバルモデルを学習結果となるモデルとして決定する。

10

【 0 0 0 5 】

連合学習では、各クライアントは、ローカルモデルまたは差分情報をサーバに提供すればよく、それぞれのクライアントが独自に保持しているデータをサーバに提供する必要がない。そして、サーバが各クライアントからデータを収集してモデルを学習した場合と同様のモデルを得ることができる。すなわち、各クライアントが独自に保持しているデータを外部に提供することなく、サーバは、モデルを得ることができる。

20

【 0 0 0 6 】

連合学習では、グローバルモデルを得ることを目的としていることが多い。これに対して、個々のクライアント毎に、個々のクライアントに適したモデルが得られるようにした技術も提案されている。このような技術は、Personalized Federated Learning と呼ばれる。一般に、それぞれのクライアントは、類似しているが異なるデータを保持する。例えば、ある地域（Aとする。）の銀行のクライアントと、別の地域（Bとする。）の銀行のクライアントとが、それぞれ、顧客の預金額のデータを学習データとして保持しているとする。この学習データは、いずれも顧客の預金額のデータであり、類似したデータである。しかし、地域性の違いによって、データの性質が異なることがある。そして、地域性の違いによって、地域Aの銀行のクライアントに適したモデルと、地域Bの銀行のクライアントに適したモデルも、異なることになる。Personalized Federated Learning では、各クライアントが、それぞれのクライアントに適したモデルを得る。

30

【 0 0 0 7 】

Personalized Federated Learning の一例が、非特許文献1に記載されている。非特許文献1に記載された技術は、FedProx と称されている。FedProx では、ローカルモデルにおける正解値と予測値のずれを評価するロス関数の出力と、グローバルモデルおよびローカルモデルのパラメータのずれとを加算した式を用いる。

【 0 0 0 8 】

40

Personalized Federated Learning の他の例が、非特許文献2に記載されている。非特許文献2に記載された技術は、FedFomo と称されている。FedFomo では、各クライアントがそれぞれ、他の各クライアントのローカルモデルを受け取り、それぞれのクライアントがそれぞれ別個に、各クライアントのローカルモデルに重み付けを行い、自身に適したモデルを得る。

【 0 0 0 9 】

また、Personalized Federated Learning とは別に、ディープラーニングに関する種々の技術も提案されている（非特許文献3，4を参照）。非特許文献3には、学習によって求めた複数の固定値を用いて、入力値に応じたそれらの固定値の重み付け和を求めることが記載されている。例えば、学習によって、 W_1 ， W_2 ， W_3 という3つの固定値を求め

50

たとする。非特許文献 3 に記載の技術 (CondConv と称されている。) では、入力値に応じて、 W_1 、 W_2 、 W_3 に対応する重み値を定め、入力値に応じた重み値で、 W_1 、 W_2 、 W_3 の重み付け和を求める。

【0010】

また、非特許文献 4 には、学習時には、並列に処理される複数の畳み込み演算のパラメータを学習し、推論時には、その複数の畳み込み演算を 1 つの畳み込み演算にまとめることが記載されている。例えば、非特許文献 4 には、学習時には、 3×3 フィルタの畳み込み演算のパラメータと、 1×1 フィルタの畳み込み演算のパラメータとを学習し、推論時には、それらの畳み込み演算を、 3×3 フィルタの 1 つの畳み込み演算にまとめることが記載されている。非特許文献 4 に記載された技術は、RepVGG と称されている。

10

【先行技術文献】

【非特許文献】

【0011】

【文献】Tian Li, et al., "Federated Optimization in Heterogeneous Networks", [2021 年 6 月 7 日検索]、インターネット、URL: <https://arxiv.org/pdf/1812.06127.pdf>

【文献】Michael Zhang, et al., "Personalized Federated Learning with First Order Model Optimization", [2021 年 6 月 7 日検索]、インターネット、URL: <https://arxiv.org/pdf/2012.08565.pdf>

【文献】Brandon Yang, et al., "CondConv: Conditionally Parameterized Convolutions for Efficient Inference", [2021 年 6 月 7 日検索]、インターネット、URL: <https://arxiv.org/pdf/1904.04971.pdf>

20

【文献】Xiaohan Ding, et al., "RepVGG: Making VGG-style ConvNets Great Again", [2021 年 6 月 7 日検索]、インターネット、URL: <https://arxiv.org/pdf/2101.03697.pdf>

【発明の概要】

【発明が解決しようとする課題】

【0012】

非特許文献 1 に記載された技術 (FedProx) では、前述のように、ローカルモデルを求めるために、ロス関数の出力と、グローバルモデルおよびローカルモデルのパラメータのずれとを加算した式を用いる。しかし、このパラメータのずれが小さくてもモデルの出力が大きく変動する場合や、このパラメータのずれが大きくてもモデルの出力があまり変動しない場合がある。すなわち、グローバルモデルとローカルモデルとのパラメータのずれは、ローカルモデルの出力の性質と関連していない。その結果、非特許文献 1 に記載された技術では、最適化が難しく、また、各クライアントにおいて、精度の高いモデルを得ることは難しい。

30

【0013】

また、非特許文献 2 に記載された技術 (FedFomo) では、個々のクライアントがそれぞれ、他の各クライアントに、自身が生成したモデルを提供する必要がある。そして、モデルから、そのモデルを学習する際に用いた学習モデルを復元する技術も存在している。そのため、個々のクライアントがそれぞれ、他の複数のクライアントに自身が生成したモデルを提供することは、データ流出を抑える観点から好ましくない。

40

【0014】

そこで、本発明は、各クライアントのデータ流出の可能性を低減することができ、各クライアントに適した精度の高いモデルのパラメータを、各クライアントそれぞれが得ることができる学習システム、学習方法、および、学習プログラム、並びに、そのようなモデルで推論を行う推論器を提供することを目的とする。

【課題を解決するための手段】

【0015】

本発明による学習システムは、サーバと、複数のクライアントとを備える学習システム

50

であって、各クライアントは、共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、重み付け和の計算に関わるパラメータとを学習する学習手段と、所定の複数の操作のパラメータと重み付け和の計算に関わるパラメータとのうち、所定の複数の操作のパラメータをサーバに送信するクライアント側パラメータ送信手段とを備え、サーバは、各クライアントから受信した所定の複数の操作のパラメータに基づいて、所定の複数の操作のパラメータを計算し直すパラメータ計算手段と、その所定の複数の操作のパラメータを各クライアントに送信するサーバ側パラメータ送信手段とを備え、所定の複数の操作の数は、複数のクライアントの数未満であることを特徴とする。

【 0 0 1 7 】

10

本発明による学習方法は、サーバと、複数のクライアントとによって行われる学習方法であって、各クライアントが、共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、重み付け和の計算に関わるパラメータとを学習し、所定の複数の操作のパラメータと重み付け和の計算に関わるパラメータとのうち、所定の複数の操作のパラメータをサーバに送信し、サーバが、各クライアントから受信した所定の複数の操作のパラメータに基づいて、所定の複数の操作のパラメータを計算し直し、その所定の複数の操作のパラメータを各クライアントに送信し、所定の複数の操作の数は、複数のクライアントの数未満であることを特徴とする。

【 発明の効果 】

20

【 0 0 1 9 】

本発明によれば、各クライアントのデータ流出の可能性を低減することができ、各クライアントに適した精度の高いモデルのパラメータを、各クライアントそれぞれが得ることができる。

【 図面の簡単な説明 】

【 0 0 2 0 】

【 図 1 】 連合学習でパラメータが学習される所定の複数の操作を示す模式図である。

【 図 2 】 所定の複数の操作 5 1 , 5 2 , 5 3 がそれぞれ複数の層を含んでいる場合を示す模式図である。

【 図 3 】 所定の複数の操作 5 1 , 5 2 , 5 3 に含まれる層の数が異なる場合を示す模式図である。

30

【 図 4 】 パラメータが学習されるモデルの例を示す模式図である。

【 図 5 】 本発明の実施形態の学習システムの構成例を示すブロック図である。

【 図 6 】 本発明の実施形態の処理経過の一例を示すフローチャートである。

【 図 7 】 本発明の実施形態の変形例における各クライアントの構成例を示すブロック図である。

【 図 8 】 変換部による変換後のモデルを示す模式図である。

【 図 9 】 本発明の実施形態の変形例における各クライアントの構成例を示すブロック図である。

【 図 1 0 】 クライアントとは別個の装置となる推論器を示すブロック図である。

40

【 図 1 1 】 本発明の実施形態やその種々の変形例におけるクライアント、サーバ、および、推論器に係るコンピュータの構成例を示す概略ブロック図である。

【 図 1 2 】 本発明の学習システムの概要を示すブロック図である。

【 発明を実施するための形態 】

【 0 0 2 1 】

以下、本発明の実施形態を図面を参照して説明する。

【 0 0 2 2 】

本発明の実施形態の学習システムは、後述するように、サーバと、複数のクライアントとを備える。本発明の実施形態では、サーバと各クライアントとが、連合学習によって、所定の複数の操作のパラメータを学習するとともに、各クライアントがそれぞれ独自に、

50

その所定の複数の操作の出力データの重み付け和の計算に関わるパラメータ（以下、単に、重み付け和の計算に関わるパラメータと記す。）を学習する。従って、所定の複数の操作のパラメータは、各クライアントで共通であるが、重み付け和の計算に関わるパラメータは、各クライアントで異なる。

【0023】

図1は、連合学習でパラメータが学習される所定の複数の操作を示す模式図である。所定の複数の操作とは、共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある複数の操作である。図1では、操作51, 52, 53が、所定の複数の操作に該当する。すなわち、操作51, 52, 53には、共通の入力データが与えられ、操作51, 52, 53それぞれの出力データの重み付け和が計算される。図1に示す α_1 , α_2 , α_3 は、出力データの重み付け和を計算する際に用いられる重み値である。各重み値 α_1 , α_2 , α_3 はそれぞれ、0以上1以下の値であり、各重み値 α_1 , α_2 , α_3 の総和は1である。

10

【0024】

図1に示す例では、所定の複数の操作51~53のパラメータが、サーバおよび各クライアントによる連合学習によって学習される。また、 α_1 , α_2 , α_3 は、重み付け和の計算に関わるパラメータであり、各クライアントそれぞれによって独自に学習される。

【0025】

また、図1では、操作51, 52, 53それぞれの出力データの重み付け和に対して正規化を行うノーマライズ操作54も図示している。ノーマライズ操作54のパラメータは、重み付け和の計算に関わるパラメータとして扱う。従って、ノーマライズ操作54のパラメータは、 α_1 , α_2 , α_3 と同様に、各クライアントそれぞれによって独自に学習される。なお、ノーマライズ操作54の一例として、ノーマライズ操作54への入力データからある数値（ β とする。）を減算し、その減算結果にある数値（ γ とする。）を乗算する処理が考えられる。この場合、 β および γ が、ノーマライズ操作54のパラメータに該当する。ただし、ノーマライズ操作54における演算やパラメータは、この例に限定されない。

20

【0026】

図1では、所定の複数の操作の数が3個である場合を示しているが、所定の複数の操作の数は3個に限定されない。ただし、所定の複数の操作の数は、複数のクライアントの数未満であるという制約がある。

30

【0027】

所定の複数の操作51, 52, 53は、複数の層を含んでもよい。図2は、所定の複数の操作51, 52, 53がそれぞれ複数の層を含んでいる場合を示す模式図である。図2では、操作51が、層A~Cを含み、操作52が、層D~Fを含み、操作53が、層G~Iを含む場合を例示している。この場合、層A~Cのパラメータが、操作51のパラメータとなる。同様に、層D~Fのパラメータが、操作52のパラメータとなり、層G~Iのパラメータが、操作53のパラメータとなる。

【0028】

図3は、所定の複数の操作51, 52, 53に含まれる層の数が異なる場合を示す模式図である。図3に示すように、所定の複数の操作51, 52, 53に含まれる層の数が、操作毎に異なってもよい。

40

【0029】

以下の説明では、説明を簡単にするために、所定の複数の操作51, 52, 53がそれぞれ、畳み込み操作である場合を例にする。畳み込み操作は、線形操作であるが、所定の複数の操作は、それぞれ、線形操作であっても、線形操作でなくてもよい。例えば、所定の複数の操作51, 52, 53が全て線形操作であってもよく、また、所定の複数の操作51, 52, 53が全て線形操作でなくてもよい。また、所定の複数の操作51, 52, 53のうちの一部が線形操作であり、残りの一部が線形操作でなくてもよい。なお、畳み込み操作以外の線形操作の例として、例えば、全結合操作が挙げられる。

50

【 0 0 3 0 】

図 4 は、パラメータが学習されるモデルの例を示す模式図である。実際のモデルは、より多くの操作が続くが、図 4 では、簡単な構成のモデルを例示している。図 4 では、畳み込み操作 5 1 , 5 2 , 5 3 は、共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある。従って、畳み込み操作 5 1 , 5 2 , 5 3 は、図 1 に示す操作 5 1 , 5 2 , 5 3 と同様に、所定の複数の操作に該当する。よって、図 1 に示す操作 5 1 , 5 2 , 5 3 と同一の符号で表す。図 2、図 3、図 4 に示す α_1 , α_2 , α_3 は、図 1 に示す α_1 , α_2 , α_3 と同様に、出力データの重み付け和を計算する際に用いられる重み値である。

【 0 0 3 1 】

畳み込み操作 5 1 のパラメータ、畳み込み操作 5 2 のパラメータ、および、畳み込み操作 5 3 のパラメータは、入力データに畳み込み演算をする際に用いる複数の重み値（以下、重み値群と記す。）である。畳み込み操作 5 1 , 5 2 , 5 3 それぞれの重み値群は、サーバと各クライアントとによって連合学習で、学習される。

【 0 0 3 2 】

ノーマライズ操作 5 4 は、畳み込み操作 5 1 , 5 2 , 5 3 それぞれの出力データの重み付け和に対して正規化を行う操作である。既に説明したように、ノーマライズ操作 5 4 のパラメータは、重み付け和の計算に関わるパラメータとして扱う。従って、ノーマライズ操作 5 4 のパラメータは、 α_1 , α_2 , α_3 と同様に、各クライアントそれぞれによって独自に学習される。

【 0 0 3 3 】

アクティベーション操作 5 5 は、ノーマライズ操作 5 4 の出力データに、活性化関数（例えば、ReLU (Rectified Linear Unit)) を適用する操作である。アクティベーション操作 5 5 はパラメータを有していなくてもよく、ここでは、活性化関数が予め定められていて、アクティベーション操作 5 5 のパラメータがない場合を例にして説明する。なお、アクティベーション操作 5 5 のパラメータが存在する場合、そのパラメータは、所定の複数の操作 5 1 , 5 2 , 5 3 のパラメータと同様に、サーバと各クライアントとによって連合学習で、学習されればよい。

【 0 0 3 4 】

図 5 は、本発明の実施形態の学習システムの構成例を示すブロック図である。以下では、図 5 に示す学習システムが、図 4 に示すモデルのパラメータを学習する場合を例にして説明する。

【 0 0 3 5 】

本発明の実施形態の学習システムは、サーバ 2 0 と、複数のクライアント 1 0_a ~ 1 0_e を備える。サーバ 2 0 と、複数のクライアント 1 0_a ~ 1 0_e とは、通信ネットワーク 3 0 を介して、通信可能に接続される。図 5 では、5 台のクライアント 1 0_a ~ 1 0_e を示したが、クライアントの台数は、5 台に限定されない。ただし、前述のように、所定の複数の操作の数は、複数のクライアントの数未満であるという制約がある。本例では、所定の複数の操作（畳み込み操作 5 1 , 5 2 , 5 3 ）の数は“ 3 ”であり（図 4 参照）、複数のクライアントの数は“ 5 ”であるので、この制約を満たしている。

【 0 0 3 6 】

各クライアント 1 0_a ~ 1 0_e は同様の構成であり、特にクライアントを区別しない場合には、クライアントを符号 1 0 で示す。

【 0 0 3 7 】

以下、図 5 を参照して、クライアント 1 0_a を例にして、クライアント 1 0 の構成を説明する。クライアント 1 0 は、学習部 1 1 と、クライアント側パラメータ送受信部 1 2 と、記憶部 1 3 とを備える。

【 0 0 3 8 】

学習部 1 1 は、機械学習により、所定の複数の操作のパラメータ（本例では、畳み込み操作 5 1 , 5 2 , 5 3 それぞれの重み値群）と、重み付け和の計算に関わるパラメータと

10

20

30

40

50

を学習する。本例では、 α_1 、 α_2 、 α_3 、および、ノーマライズ操作 5 4 のパラメータが、重み付け和の計算に関わるパラメータに該当する。

【0039】

記憶部 1 3 は、学習部 1 1 が上記の各種パラメータを学習する際に用いる学習データ、および、学習されたパラメータによって定まるモデルを記憶する記憶装置である。

【0040】

各クライアント 1 0_a ~ 1 0_e の記憶部 1 3 には、それぞれ、クライアント独自の学習データが予め記憶されている。

【0041】

クライアント側パラメータ送受信部 1 2 は、サーバ 2 0 に、所定の複数の操作のパラメータ（本例では、畳み込み操作 5 1、5 2、5 3 それぞれの重み値群）と重み付け和の計算に関わるパラメータ（本例では、 α_1 、 α_2 、 α_3 、および、ノーマライズ操作 5 4 のパラメータ）とのうち、所定の複数の操作のパラメータを送信する。

10

【0042】

従って、重み付け和の計算に関わるパラメータ（ α_1 、 α_2 、 α_3 、および、ノーマライズ操作 5 4 のパラメータ）は、サーバ 2 0 に送信されない。このことは、重み付け和の計算に関わるパラメータは、連合学習によって学習されずに、それぞれのクライアント 1 0_a ~ 1 0_e の学習部 1 1 が独自に、重み付け和の計算に関わるパラメータを学習することを意味する。

【0043】

また、クライアント側パラメータ送受信部 1 2 は、サーバ 2 0 において計算し直された所定の複数の操作のパラメータ（畳み込み操作 5 1、5 2、5 3 それぞれの重み値群）を、サーバ 2 0 から受信する。

20

【0044】

各クライアント 1 0 はそれぞれ、例えば、コンピュータによって実現される。そして、クライアント側パラメータ送受信部 1 2 は、例えば、学習プログラムに従って動作する CPU（Central Processing Unit）、および、そのコンピュータの通信インタフェースによって実現される。例えば、CPU が、コンピュータのプログラム記憶装置等のプログラム記録媒体から学習プログラムを読み込み、その学習プログラムに従って、通信インタフェースを用いて、クライアント側パラメータ送受信部 1 2 として動作すればよい。なお、通信インタフェースは、通信ネットワーク 3 0 とのインタフェースである。また、学習部 1 1 は、例えば、学習プログラムに従って動作する CPU によって実現される。例えば、CPU が上記のようにプログラム記録媒体から学習プログラムを読み込み、その学習プログラムに従って、学習部 1 1 として動作すればよい。

30

【0045】

サーバ 2 0 は、パラメータ計算部 2 1 と、サーバ側パラメータ送受信部 2 2 とを備える。

【0046】

サーバ側パラメータ送受信部 2 2 は、各クライアント 1 0 のクライアント側パラメータ送受信部 1 2 が送信した所定の複数の操作のパラメータ（畳み込み操作 5 1、5 2、5 3 それぞれの重み値群）を、各クライアント 1 0 から受信する。

40

【0047】

また、サーバ側パラメータ送受信部 2 2 は、パラメータ計算部 2 1 によって計算し直された所定の複数の操作のパラメータ（畳み込み操作 5 1、5 2、5 3 それぞれの重み値群）を各クライアント 1 0 に送信する。この所定の複数の操作のパラメータは、各クライアント 1 0 のクライアント側パラメータ送受信部 1 2 によって受信される。

【0048】

パラメータ計算部 2 1 は、サーバ側パラメータ送受信部 2 2 が各クライアント 1 0 から受信した所定の複数の操作のパラメータ（畳み込み操作 5 1、5 2、5 3 それぞれの重み値群）に基づいて、所定の複数の操作のパラメータを計算し直す。

【0049】

50

例えば、畳み込み操作 5 1 の重み値群に属する重み値は、クライアント 1 0_a ~ 1 0_e の違いにより、クライアント 1 0 毎にそれぞれ異なる。ただし、畳み込み操作 5 1 の重み値群に属する個々の重み値は、各クライアント 1 0_a ~ 1 0_e で対応している。パラメータ計算部 2 1 は、畳み込み操作 5 1 の重み値群に属する重み値毎に、クライアント 1 0_a で得られた重み値、クライアント 1 0_b で得られた重み値、クライアント 1 0_c で得られた重み値、クライアント 1 0_d で得られた重み値、クライアント 1 0_e で得られた重み値の平均値を計算することによって、畳み込み操作 5 1 の重み値群を計算し直す。同様に、パラメータ計算部 2 1 は、畳み込み操作 5 2 の重み値群を計算し直す。また、同様に、パラメータ計算部 2 1 は、畳み込み操作 5 3 の重み値群を計算し直す。

【 0 0 5 0 】

10

前述のように、サーバ側パラメータ送受信部 2 2 は、パラメータ計算部 2 1 によって計算し直された所定の複数の操作のパラメータ（畳み込み操作 5 1 , 5 2 , 5 3 それぞれの重み値群）を各クライアント 1 0 に送信する。

【 0 0 5 1 】

なお、各クライアント 1 0 の学習部 1 1 は、それぞれ、独自に保持する学習データと、サーバ 2 0 から受信した所定の複数の操作のパラメータとを用いて、再度、機械学習により、所定の複数の操作のパラメータを学習し、また、重み付け和の計算に関わるパラメータを学習する。

【 0 0 5 2 】

サーバ 2 0 は、例えば、コンピュータによって実現される。そして、サーバ側パラメータ送受信部 2 2 は、例えば、サーバ用プログラムに従って動作する CPU、および、そのコンピュータの通信インタフェースによって実現される。例えば、CPU が、コンピュータのプログラム記憶装置等のプログラム記録媒体からサーバ用プログラムを読み込み、そのサーバ用プログラムに従って、通信インタフェースを用いて、サーバ側パラメータ送受信部 2 2 として動作すればよい。通信インタフェースは、通信ネットワーク 3 0 とのインタフェースである。また、パラメータ計算部 2 1 は、例えば、サーバ用プログラムに従って動作する CPU によって実現される。例えば、CPU が上記のようにプログラム記録媒体からサーバ用プログラムを読み込み、そのサーバ用プログラムに従って、パラメータ計算部 2 1 として動作すればよい。

20

【 0 0 5 3 】

30

次に、本発明の実施形態の処理経過について説明する。図 6 は、本発明の実施形態の処理経過の一例を示すフローチャートである。なお、図 6 は一例であり、本発明の実施形態の処理経過は、図 6 に示す例に限定されない。

【 0 0 5 4 】

また、図 6 では、サーバ 2 0 およびクライアント 1 0_a の動作を図示しているが、クライアント 1 0_b ~ 1 0_e の動作は、クライアント 1 0_a の動作と同様である。ただし、各クライアント 1 0 が記憶部 1 3 に保持している学習データは、クライアント 1 0 毎に異なる。

【 0 0 5 5 】

クライアント 1 0_a の学習部 1 1 は、記憶部 1 3 に記憶されている学習データに基づいて、機械学習により、所定の複数の操作のパラメータ（畳み込み操作 5 1 , 5 2 , 5 3 それぞれの重み値群）を学習し、また、重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作 5 4 のパラメータ）を学習する（ステップ S 1 ）。

40

【 0 0 5 6 】

他の各クライアント 1 0_b ~ 1 0_e の学習部 1 1 も同様に、所定の複数の操作のパラメータを学習し、また、重み付け和の計算に関わるパラメータを学習する。

【 0 0 5 7 】

次に、クライアント 1 0_a のクライアント側パラメータ送受信部 1 2 は、ステップ S 1 で学習した所定の複数の操作のパラメータ（畳み込み操作 5 1 , 5 2 , 5 3 それぞれの重み値群）と重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作 5 4 のパラメータ）とのうち、所定の複数の操作のパラメータをサーバ 2 0 に送信

50

する（ステップS2）。

【0058】

他の各クライアント10_b～10_eのクライアント側パラメータ送受信部12もそれぞれ、同様に、所定の複数の操作のパラメータと重み付け和の計算に関わるパラメータとのうち、所定の複数の操作のパラメータをサーバ20に送信する。

【0059】

従って、重み付け和の計算に関わるパラメータ（ α_1 、 α_2 、 α_3 、および、ノーマライズ操作54のパラメータ）は、各クライアント10_a～10_eからサーバ20に送信されない。

【0060】

サーバ20のサーバ側パラメータ送受信部22は、各クライアント10_a～10_eから、所定の複数の操作のパラメータ（畳み込み操作51、52、53それぞれの重み値群）を受信する。

【0061】

そして、サーバ20のパラメータ計算部21は、各クライアント10_a～10_eから受信した所定の複数の操作のパラメータに基づいて、所定の複数のパラメータを計算し直す（ステップS3）。パラメータ計算部21が所定の複数のパラメータを計算し直す動作の例については、既に説明したので、ここでは説明を省略する。

【0062】

次に、サーバ側パラメータ送受信部22は、ステップS3で計算し直された所定の複数の操作のパラメータ（畳み込み操作51、52、53それぞれの重み値群）を、各クライアント10_a～10_eに送信する（ステップS4）。ステップS4では、各クライアント10_a～10_eに同一のパラメータが送信される。

【0063】

ステップS4で送信されたパラメータを受信した各クライアント10_a～10_eは、ステップS1以降の処理を繰り返す。ただし、サーバ20で計算し直された所定の複数の操作のパラメータを受信した後にステップS1を行う場合、クライアント10_aの学習部11は、その所定の複数の操作のパラメータ、および、記憶部13に記憶されている学習データに基づいて、機械学習により、所定の複数の操作のパラメータ（畳み込み操作51、52、53それぞれの重み値群）を学習し、また、重み付け和の計算に関わるパラメータ（ α_1 、 α_2 、 α_3 、および、ノーマライズ操作54のパラメータ）を学習する。他のクライアント10_b～10_eの学習部11も同様である。

【0064】

各クライアント10_a～10_eがステップS1以降の処理を繰り返すことによって、各クライアント10およびサーバ20において、ステップS1～S4の処理が繰り返される。例えば、ステップS1～S4の処理の繰り返し数が所定回数に達したことが、各クライアント10およびサーバ20による学習（換言すれば、連合学習）の終了条件であると予め定められていてもよい。この場合、例えば、各クライアント10の学習部11は、ステップS1の実行回数をカウントし、ステップS1の実行回数が所定回数に達したときにおける、所定の複数の操作のパラメータ（畳み込み操作51、52、53それぞれの重み値群）、および、重み付け和の計算に関わるパラメータ（ α_1 、 α_2 、 α_3 、および、ノーマライズ操作54のパラメータ）を、それぞれのパラメータの確定値であると決定し、それらのパラメータによって定まるモデルを記憶部13に記憶させてもよい。

【0065】

なお、各クライアント10およびサーバ20による学習の終了条件は、上記の例に限定されず、他の条件であってもよい。

【0066】

本実施形態によれば、所定の複数の操作のパラメータ（畳み込み操作51、52、53それぞれの重み値群）は、各クライアント10およびサーバ20による学習（連合学習）によって定まる。一方、重み付け和の計算に関わるパラメータ（ α_1 、 α_2 、 α_3 、および、

10

20

30

40

50

、ノーマライズ操作 5 4 のパラメータ) は、各クライアント 1 0 の学習部 1 1 が独自に学習する。所定の複数の操作のパラメータを各クライアント 1 0 で共通のパラメータとしつつ、各クライアント独自のパラメータを、各クライアント 1 0 は得ることができる。すなわち、共通なパラメータを含みつつも、クライアント 1 0 で個別のパラメータが得られる。そして、本実施形態では、FedProx (非特許文献 1 参照) とは異なり、モデルの性質と関連のないパラメータのずれ (グローバルモデルとローカルモデルとのパラメータのずれ) は用いていない。よって、各クライアント 1 0 は、それぞれのクライアント 1 0 に適したパラメータを得ることができ、そのパラメータによって定まる精度の高いモデルを得ることができる。

【 0 0 6 7 】

10

さらに、本実施形態では、各クライアント 1 0 は、サーバ 2 0 とパラメータを授受するが、他の各クライアントとの間でモデルの授受は行っていない。従って、FedFomo (非特許文献 2 参照) よりも、データ流出の可能性を低減することができる。

【 0 0 6 8 】

また、所定の複数の操作の数は、複数のクライアントの数未満である。従って、所定の複数の操作のうち、クライアントにおいて重要となる操作は、一部のクライアントで共通する。例えば、 α_1 の値が大きくなるという事象は、一部のクライアントで共通する。同様に、 α_2 の値が大きくなるという事象も、一部のクライアントで共通し、 α_3 の値が大きくなるという事象も、一部のクライアントで共通する。この結果、各クライアント 1 0 に適したパラメータが得られ、そのパラメータにより各クライアントに適したモデルが得られる。さらに、そのモデルの性質が互いに大きく異なることになることを防止できる。

20

【 0 0 6 9 】

所定の複数の操作の数が、複数のクライアントの数よりも大きい場合を考える。例えば、所定の複数の操作の数が 6 個であり、クライアントの台数が 3 台であるとする。この場合、各操作に対する重み値である $\alpha_1 \sim \alpha_6$ がパラメータとなる。このとき、1 台目のクライアントでは、 α_1, α_2 が大きくなり、2 台目のクライアントでは、 α_3, α_4 が大きくなり、3 台目のクライアントでは、 α_5, α_6 が大きくなるということが生じ得る。この場合、クライアントにおいて重要となる操作は、3 台のクライアントにおいて異なることになり、3 台のクライアントのモデルの性質が大きく異なることになってしまう。所定の複数の操作の数が、複数のクライアントの数未満であることによって、このようなことを防止することができる。すなわち、各クライアントのモデルの性質が離れすぎることを防止できる。よって、個々のクライアントに適したモデルを得ることができ、また、各クライアントのモデルの性質が大きく異なることを防ぐことができる。

30

【 0 0 7 0 】

次に、本実施形態の変形例を説明する。図 6 に示すフローチャートでは、学習部 1 1 が、ステップ S 1 で、所定の複数の操作のパラメータを学習し、また、重み付け和の計算に関わるパラメータを学習する場合を示した。ステップ S 1 で、各クライアント 1 0 の学習部 1 1 は、所定の複数の操作のパラメータを学習し、重み付け和の計算に関わるパラメータに関しては学習しなくてもよい。この場合、各クライアント 1 0 の学習部 1 1 は、所定の複数の操作のパラメータが確定した後に、それぞれ独自に、重み付け和の計算に関わるパラメータを学習すればよい。

40

【 0 0 7 1 】

次に、本実施形態の他の変形例を説明する。図 7 は、本変形例における各クライアントの構成例を示すブロック図である。なお、上記の実施形態と同様の要素には、図 5 と同一の符号を付し、説明を省略する。また、サーバ 2 0 の構成および動作は、上記の実施形態のサーバ 2 0 の構成および動作と同様であり、説明を省略する。

【 0 0 7 2 】

本変形例では、所定の複数の操作が、全て線形操作であるものとする。そのため、本変形例でも、図 4 を参照して説明する。ただし、所定の複数の操作は、全て線形操作であればよく、図 4 に示すように、所定の複数の操作が全て畳み込み操作である場合に限定され

50

ない。

【0073】

本変形例では、クライアント10は、学習部11、クライアント側パラメータ送受信部12、記憶部13に加えて、変換部14を備える。

【0074】

所定の複数の操作のパラメータ（畳み込み操作51, 52, 53それぞれの重み値群）、および、重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作54のパラメータ）の確定値が決定され、それらのパラメータによって定まるモデルが記憶部13に記憶されるまでの動作は、上記の実施形態と同様である。

【0075】

変換部14は、そのように所定の複数の操作のパラメータ（畳み込み操作51, 52, 53それぞれの重み値群）、および、重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作54のパラメータ）が確定した後に、所定の複数の操作のパラメータ、および、重み付け和の計算に関わるパラメータに基づいて、その所定の複数の操作のパラメータを1つの操作に変換する。

【0076】

図4に示す例では、変換部14は、畳み込み操作51, 52, 53それぞれの重み値群、並びに、 α_1 , α_2 , α_3 、および、ノーマライズ操作54のパラメータが確定した後に、畳み込み操作51, 52, 53それぞれの重み値群、および、 α_1 , α_2 , α_3 に基づいて、畳み込み操作51, 52, 53を1つの畳み込み操作に変換する。

【0077】

入力データは、複数の数値を含むが、ここでは、便宜的に1つの符号 x で表す。畳み込み操作51の重み値群も複数の重み値を含むが、ここでは、便宜的に1つの符号 w_1 で表す。同様に、畳み込み操作52の重み値群、畳み込み操作53の重み値群もそれぞれ、便宜的に符号 w_2 , w_3 で表す。

【0078】

入力データ x に対する畳み込み操作51によって得られる出力データを、 $w_1 * x$ と表すこととする。同様に、入力データ x に対する畳み込み操作52によって得られる出力データを、 $w_2 * x$ と表すこととする。同様に、入力データ x に対する畳み込み操作53によって得られる出力データを、 $w_3 * x$ と表すこととする。

【0079】

この場合、出力データの重み付け和は、 $\alpha_1 (w_1 * x) + \alpha_2 (w_2 * x) + \alpha_3 (w_3 * x)$ となる。畳み込み操作51, 52, 53は線形操作であるので、この重み付け和は、 $(\alpha_1 w_1 + \alpha_2 w_2 + \alpha_3 w_3) * x$ に変形できる。従って、変換部14は、3つの畳み込み操作51, 52, 53を、 $(\alpha_1 w_1 + \alpha_2 w_2 + \alpha_3 w_3)$ を重み値群として持つ1つの畳み込み操作に変換する。図8は、変換部14による変換後のモデルを示す模式図である。図8に示す1つの畳み込み操作50は、3つの畳み込み操作51, 52, 53から、畳み込み操作51, 52, 53それぞれの重み値群、および、 α_1 , α_2 , α_3 に基づいて変換された操作である。畳み込み操作50の重み値群（パラメータ）は、上記のように、模式的に $(\alpha_1 w_1 + \alpha_2 w_2 + \alpha_3 w_3)$ と表すことができる。

【0080】

変換部14は、変換後の操作およびその操作のパラメータによって定まるモデルを、記憶部13に記憶させる。

【0081】

ここでは、所定の複数の操作が3個である場合を例にしたが、所定の複数の操作が2個または4個以上であっても、変換部14は、所定の複数の操作を1つの操作に変換することができる。また、ここでは、所定の複数の操作が全て畳み込み操作である場合を例にしたが、所定の複数の操作が全て線形操作であれば、変換部14は、所定の複数の操作を1つの操作に変換することができる。

【0082】

10

20

30

40

50

本変形例によれば、所定の複数の操作を1つの操作に変換することによって、モデルが簡略化される。従って、モデルに基づいて推論を行う場合に、計算量を減少させることができる。例えば、図4と図8とを比較した場合、図4に示すモデルでは、推論時に、3回の畳み込み操作を行わなければならない。一方、図8に示すモデルでは、推論時に行う畳み込み操作の回数は1回でよい。

【0083】

変換部14は、例えば、学習プログラムに従って動作するコンピュータのCPUによって実現される。例えば、CPUが、コンピュータのプログラム記憶装置等のプログラム記録媒体から学習プログラムを読み込み、その学習プログラムに従って、変換部14として動作すればよい。

10

【0084】

次に、本実施形態の他の変形例を説明する。図9は、本変形例における各クライアントの構成例を示すブロック図である。なお、上記の実施形態と同様の要素には、図5と同一の符号を付し、説明を省略する。また、サーバ20の構成および動作は、上記の実施形態のサーバ20の構成および動作と同様であり、説明を省略する。

【0085】

本変形例では、クライアント10は、学習部11、クライアント側パラメータ送受信部12、記憶部13に加えて、推論部15を備える。

【0086】

所定の複数の操作のパラメータ（畳み込み操作51, 52, 53それぞれの重み値群）、および、重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作54のパラメータ）の確定値が決定され、それらのパラメータによって定まるモデルが記憶部13に記憶されるまでの動作は、上記の実施形態と同様である。

20

【0087】

そのように所定の複数の操作のパラメータ（畳み込み操作51, 52, 53それぞれの重み値群）、および、重み付け和の計算に関わるパラメータ（ α_1 , α_2 , α_3 、および、ノーマライズ操作54のパラメータ）が確定し、所定の複数の操作のパラメータ、および、重み付け和の計算に関わるパラメータによって定まるモデルが記憶部13に記憶された場合、推論部15は、そのモデルに基づいて推論を行う。

【0088】

推論部15には、入力インタフェース（図示略）を介して、データが入力される。推論部15は、そのデータをモデルにおける最初の操作の入力データとし、その操作の出力データを計算する。そして、推論部15は、その出力データをモデルにおける次の操作の入力データとし、その操作の出力データを計算する。推論部15は、この動作を、モデルの最後の操作まで繰り返し、最後の操作の出力データを、推論結果として導出する。推論部15は、推論部15に入力されたデータとモデルとに基づいて得られた推論結果を、例えば、クライアント10が備えるディスプレイ装置（図示略）に表示してもよい。

30

【0089】

本変形例によれば、パラメータが確定したことによって定めるモデルを得ることができるだけでなく、そのモデルを用いて、推論を行うことができる。

40

【0090】

推論部15は、例えば、学習プログラムに従って動作するコンピュータのCPUによって実現される。例えば、CPUが、コンピュータのプログラム記憶装置等のプログラム記録媒体から学習プログラムを読み込み、その学習プログラムに従って、推論部15として動作すればよい。

【0091】

本変形例のクライアント10は、モデルに基づいて推論を行う推論器であるということができる。

【0092】

また、クライアント10とは別個の装置として推論器が設けられていてもよい。図10

50

は、クライアント１０とは別個の装置となる推論器を示すブロック図である。図１０に示す推論器４０は、記憶部４１と、推論部１５とを備える。

【００９３】

記憶部４１は、上記の実施形態またはその種々の変形例において、クライアント１０の記憶部１３に記憶されたモデルと同一のモデルを記憶する記憶装置である。上記の実施形態またはその種々の変形例におけるクライアント１０の記憶部１３に記憶されたモデルを、推論器４０の記憶部４１にコピーして、記憶部４１にモデルを記憶させておけばよい。

【００９４】

推論部１５は、図９に示すクライアント１０が備える推論部１５と同様である。すなわち、推論部１５には、入力インタフェース（図示略）を介して、データが入力される。推論部１５は、そのデータをモデルにおける最初の操作の入力データとし、その操作の出力データを計算する。そして、推論部１５は、その出力データをモデルにおける次の操作の入力データとし、その操作の出力データを計算する。推論部１５は、この動作を、モデルの最後の操作まで繰り返し、最後の操作の出力データを、推論結果として導出する。推論部１５は、その推論結果を、例えば、推論器４０が備えるディスプレイ装置（図示略）に表示してもよい。

【００９５】

推論器４０は、例えば、コンピュータによって実現され、推論部１５は、例えば、推論プログラムに従って動作するそのコンピュータのＣＰＵによって実現される。

【００９６】

上記の種々の変形例は、組合せて実現されてもよい。例えば、クライアント１０が変換部１４（図７参照）と、推論部１５（図９参照）とを備えていてもよい。

【００９７】

また、上記の実施形態やその種々の変形例では、図４に示す簡単な構成のモデルを例にして説明した。本発明の実施形態やその種々の変形例が学習の対象とするモデルは、所定の複数の操作を、複数の箇所に含むモデルであってもよい。

【００９８】

そして、所定の複数の操作が、モデル内の複数の箇所に存在する場合、その箇所毎に、所定の複数の操作の数が異なってもよく、あるいは、各箇所で、所定の複数の操作の数が同一であってもよい。そして、各箇所で、所定の複数の操作の数が同一である場合、出力データの重み付け和を計算する際に用いられる重み値の数も各箇所で同一である。このとき、所定の複数の操作の数を n とした場合、各操作に対応する重み値は、 $\alpha_1, \dots, \alpha_n$ と表すことができる。そして、各箇所における α_i （ i は、１以上 n 以下の整数）を共通の値としてもよい。例えば、学習部１１は、各箇所における α_1 を、共通の値として学習してもよい。 $\alpha_2 \sim \alpha_n$ に関しても同様である。

【００９９】

図１１は、本発明の実施形態やその種々の変形例におけるクライアント１０、サーバ２０、および、推論器４０に係るコンピュータの構成例を示す概略ブロック図である。以下、図１１を参照して説明するが、クライアント１０として用いられるコンピュータと、サーバ２０として用いられるコンピュータと、推論器４０として用いられるコンピュータとは、別々のコンピュータである。

【０１００】

コンピュータ１０００は、ＣＰＵ１００１と、主記憶装置１００２と、補助記憶装置１００３と、インタフェース１００４と、通信インタフェース１００５とを備える。

【０１０１】

本発明の実施形態やその種々の変形例におけるクライアント１０、サーバ２０、および、推論器４０は、例えば、コンピュータ１０００によって実現される。ただし、前述のように、クライアント１０として用いられるコンピュータと、サーバ２０として用いられるコンピュータと、推論器４０として用いられるコンピュータとは、別々のコンピュータである。

10

20

30

40

50

【 0 1 0 2 】

クライアント 1 0 として用いられるコンピュータ 1 0 0 0 の動作は、学習プログラムの形式で補助記憶装置 1 0 0 3 に記憶されている。CPU 1 0 0 1 は、その学習プログラムを補助記憶装置 1 0 0 3 から読み出して、主記憶装置 1 0 0 2 に展開し、その学習プログラムに従って、上記の実施形態やその種々の変形例におけるクライアント 1 0 として動作する。なお、クライアント 1 0 として用いられるコンピュータ 1 0 0 0 は、ディスプレイ装置や、データが入力される入力インタフェースを備えていてもよい。

【 0 1 0 3 】

サーバ 2 0 として用いられるコンピュータ 1 0 0 0 の動作は、サーバ用プログラムの形式で補助記憶装置 1 0 0 3 に記憶されている。CPU 1 0 0 1 は、そのサーバ用プログラムを補助記憶装置 1 0 0 3 から読み出して、主記憶装置 1 0 0 2 に展開し、そのサーバ用プログラムに従って、上記の実施形態やその種々の変形例におけるサーバ 2 0 として動作する。

【 0 1 0 4 】

図 1 0 に示す推論器 4 0 として用いられるコンピュータ 1 0 0 0 の動作は、推論プログラムの形式で補助記憶装置 1 0 0 3 に記憶されている。CPU 1 0 0 1 は、その推論プログラムを補助記憶装置 1 0 0 3 から読み出して、主記憶装置 1 0 0 2 に展開し、その推論プログラムに従って、推論器 4 0 として動作する。なお、推論器 4 0 として用いられるコンピュータ 1 0 0 0 は、通信インタフェース 1 0 0 5 を備えていなくてもよい。また、推論器 4 0 として用いられるコンピュータ 1 0 0 0 は、ディスプレイ装置や、データが入力される入力インタフェースを備えていてもよい。

【 0 1 0 5 】

補助記憶装置 1 0 0 3 は、一時的でない有形の媒体の例である。一時的でない有形の媒体の他の例として、インタフェース 1 0 0 4 を介して接続される磁気ディスク、光磁気ディスク、CD - ROM (Compact Disk Read Only Memory)、DVD - ROM (Digital Versatile Disk Read Only Memory)、半導体メモリ等が挙げられる。また、プログラムが通信回線によってコンピュータ 1 0 0 0 に配信される場合、配信を受けたコンピュータ 1 0 0 0 がそのプログラムを主記憶装置 1 0 0 2 に展開し、そのプログラムに従って動作してもよい。

【 0 1 0 6 】

また、クライアント 1 0 の各構成要素の一部または全部は、汎用または専用の回路 (circuitry)、プロセッサ等やこれらの組み合わせによって実現されてもよい。これらは、単一のチップによって構成されてもよいし、バスを介して接続される複数のチップによって構成されてもよい。各構成要素の一部または全部は、上述の回路等とプログラムとの組み合わせによって実現されてもよい。この点は、サーバ 2 0 や、図 1 0 に示す推論器 4 0 に関しても同様である。

【 0 1 0 7 】

次に、本発明の概要について説明する。図 1 2 は、本発明の学習システムの概要を示すブロック図である。

【 0 1 0 8 】

本発明の学習システムは、サーバ 1 2 0 (例えば、サーバ 2 0) と、複数のクライアント 1 1 0 (例えば、クライアント 1 0) とを備える。

【 0 1 0 9 】

各クライアント 1 1 0 は、学習手段 1 1 1 (例えば、学習部 1 1) と、クライアント側パラメータ送信手段 1 1 2 (例えば、クライアント側パラメータ送受信部 1 2) とを備える。

【 0 1 1 0 】

学習手段 1 1 1 は、共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作 (例えば、操作 5 1 , 5 2 , 5 3) のパラメータと、重み付け和の計算に関わるパラメータ (例えば、 α_1 , α_2

10

20

30

40

50

, α_3 、および、ノーマライズ操作 5 4 のパラメータ)とを学習する。

【0111】

クライアント側パラメータ送信手段 1 1 2 は、所定の複数の操作のパラメータと重み付け和の計算に関わるパラメータとのうち、所定の複数の操作のパラメータをサーバ 1 2 0 に送信する。

【0112】

サーバ 1 2 0 は、パラメータ計算手段 1 2 1 (例えば、パラメータ計算部 2 1)と、サーバ側パラメータ送信手段 1 2 2 (例えば、サーバ側パラメータ送受信部 2 2)とを備える。

【0113】

パラメータ計算手段 1 2 1 は、各クライアントから受信した所定の複数の操作のパラメータに基づいて、所定の複数の操作のパラメータを計算し直す。

【0114】

サーバ側パラメータ送信手段 1 2 2 は、その所定の複数の操作のパラメータを各クライアント 1 1 0 に送信する。

【0115】

そのような構成によって、各クライアントのデータ流出の可能性を低減することができ、各クライアントに適した精度の高いモデルのパラメータを、各クライアントそれぞれが得ることができる。

【0116】

上記の本発明の実施形態およびその変形例は、以下の付記のようにも記載され得るが、以下に限定されるわけではない。

【0117】

(付記 1)

サーバと、複数のクライアントとを備える学習システムであって、
各クライアントは、
共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、前記重み付け和の計算に関わるパラメータとを学習する学習手段と、
前記所定の複数の操作のパラメータと前記重み付け和の計算に関わるパラメータとのうち、前記所定の複数の操作のパラメータを前記サーバに送信するクライアント側パラメータ送信手段とを備え、
前記サーバは、
前記各クライアントから受信した前記所定の複数の操作のパラメータに基づいて、前記所定の複数の操作のパラメータを計算し直すパラメータ計算手段と、
前記所定の複数の操作のパラメータを前記各クライアントに送信するサーバ側パラメータ送信手段とを備える
ことを特徴とする学習システム。

【0118】

(付記 2)

前記各クライアントの前記学習手段は、前記重み付け和の計算に関わるパラメータを独自に学習する
付記 1 に記載の学習システム。

【0119】

(付記 3)

前記所定の複数の操作の数は、前記複数のクライアントの数未満である
付記 1 または付記 2 に記載の学習システム。

【0120】

(付記 4)

前記所定の複数の操作は、いずれも線形操作である

10

20

30

40

50

付記 1 から付記 3 のうちのいずれかに記載の学習システム。

【 0 1 2 1 】

(付記 5)

前記各クライアントは、

前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータが確定した場合に、前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータに基づいて、前記所定の複数の操作を 1 つの操作に変換する変換手段を備える付記 4 に記載の学習システム。

【 0 1 2 2 】

(付記 6)

前記各クライアントは、

前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータが確定した場合に、前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータによって定まるモデルに基づいて、与えられたデータに対する推論結果を導出する推論手段を備える

付記 1 から付記 5 のうちのいずれかに記載の学習システム。

【 0 1 2 3 】

(付記 7)

付記 1 から付記 6 のうちのいずれかに記載の学習システムによって得られた前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータによって定まるモデルに基づいて、与えられたデータに対する推論結果を導出する推論手段を備えることを特徴とする推論器。

【 0 1 2 4 】

(付記 8)

サーバと、複数のクライアントとによって行われる学習方法であって、
各クライアントが、

共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、前記重み付け和の計算に関わるパラメータとを学習し、

前記所定の複数の操作のパラメータと前記重み付け和の計算に関わるパラメータとのうち、前記所定の複数の操作のパラメータを前記サーバに送信し、

前記サーバが、

前記各クライアントから受信した前記所定の複数の操作のパラメータに基づいて、前記所定の複数の操作のパラメータを計算し直し、

前記所定の複数の操作のパラメータを前記各クライアントに送信する

ことを特徴とする学習方法。

【 0 1 2 5 】

(付記 9)

前記各クライアントが、前記重み付け和の計算に関わるパラメータを独自に学習する
付記 8 に記載の学習方法。

【 0 1 2 6 】

(付記 1 0)

前記所定の複数の操作の数は、前記複数のクライアントの数未満である

付記 8 または付記 9 に記載の学習方法。

【 0 1 2 7 】

(付記 1 1)

前記所定の複数の操作は、いずれも線形操作である

付記 8 から付記 1 0 のうちのいずれかに記載の学習方法。

【 0 1 2 8 】

(付記 1 2)

10

20

30

40

50

前記各クライアントが、

前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータが確定した場合に、前記所定の複数の操作のパラメータ、および、前記重み付け和の計算に関わるパラメータに基づいて、前記所定の複数の操作を１つの操作に変換する

付記１１に記載の学習方法。

【０１２９】

（付記１３）

コンピュータに、

共通の入力データが与えられるという関係にあり、かつ、出力データの重み付け和が計算されるという関係にある所定の複数の操作のパラメータと、前記重み付け和の計算に関わるパラメータとを学習する学習処理、および、

10

前記所定の複数の操作のパラメータと前記重み付け和の計算に関わるパラメータとのうち、前記所定の複数の操作のパラメータを前記サーバに送信するパラメータ送信処理を実行させるための学習プログラムを記録したコンピュータ読取可能な記録媒体。

【０１３０】

以上、実施形態を参照して本願発明を説明したが、本願発明は上記の実施形態に限定されるものではない。本願発明の構成や詳細には、本願発明の範囲内で当業者が理解し得る様々な変更をすることができる。

【産業上の利用の可能性】

【０１３１】

20

本発明は、モデルのパラメータを学習する学習システムに好適に適用可能である。

【符号の説明】

【０１３２】

１０ クライアント

１１ 学習部

１２ クライアント側パラメータ送受信部

１３ 記憶部

１４ 変換部

１５ 推論部

２０ サーバ

30

２１ パラメータ計算部

２２ サーバ側パラメータ送受信部

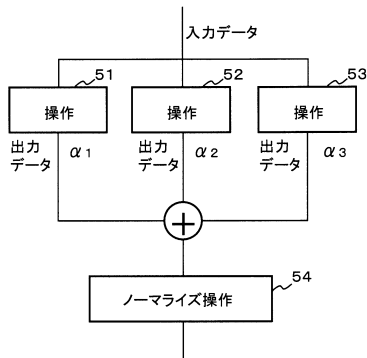
４０ 推論器

40

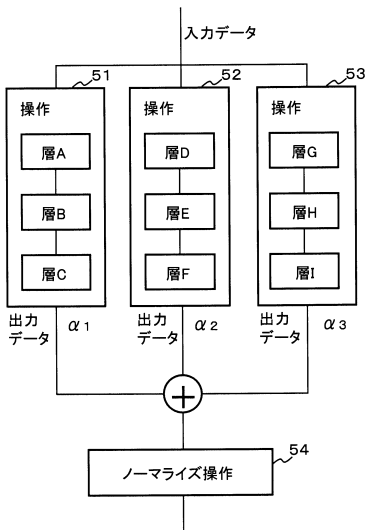
50

【図面】

【図 1】

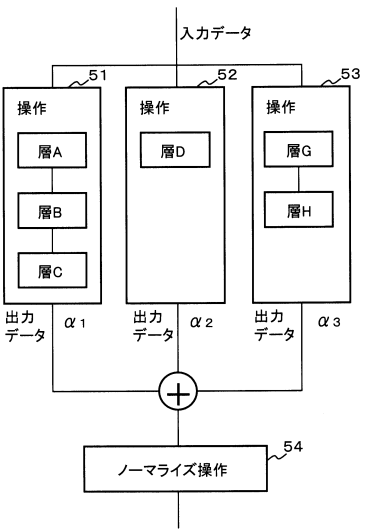


【図 2】

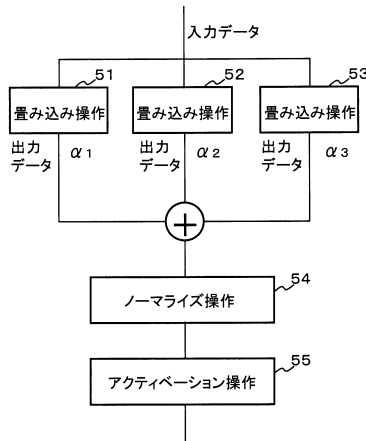


10

【図 3】



【図 4】



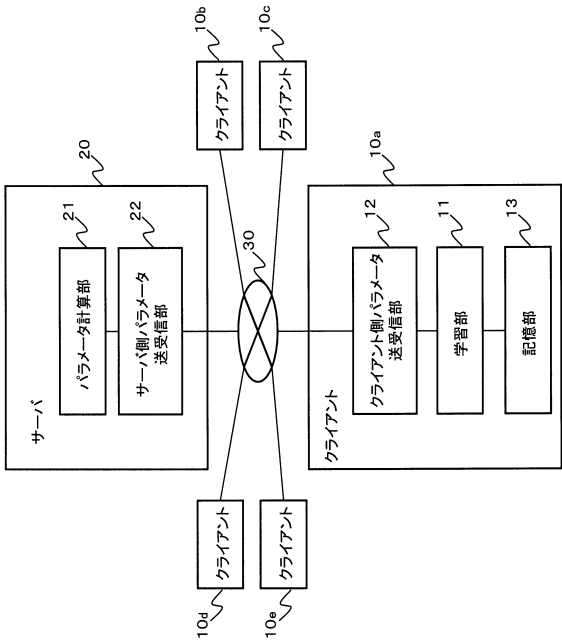
20

30

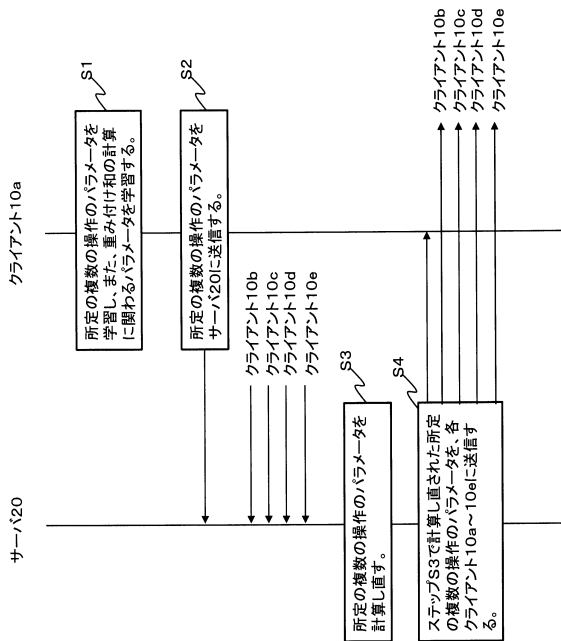
40

50

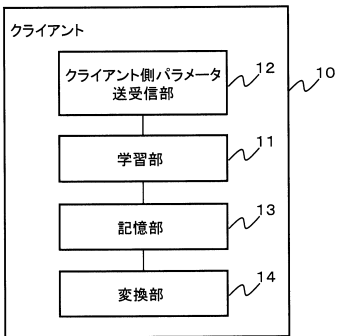
【図 5】



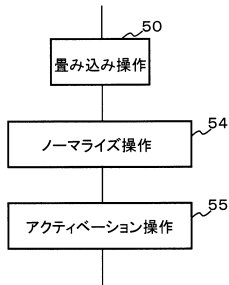
【図 6】



【図 7】



【図 8】



10

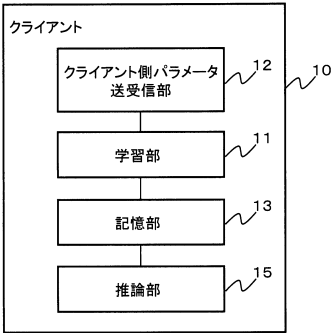
20

30

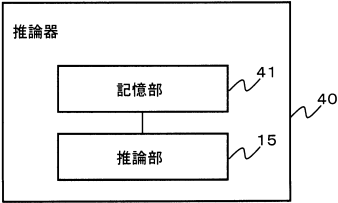
40

50

【図 9】

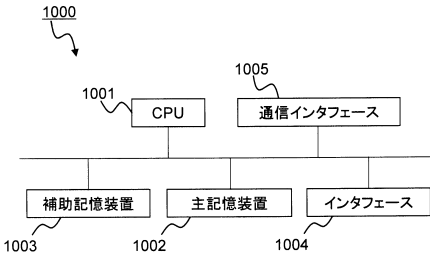


【図 10】

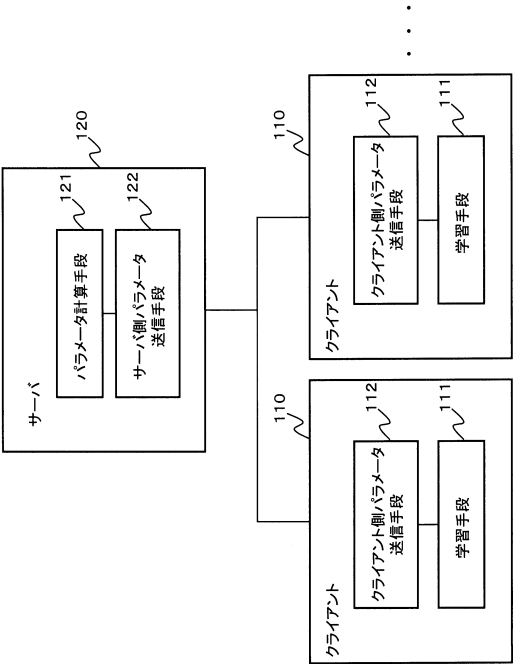


10

【図 11】



【図 12】



20

30

40

50

フロントページの続き

(56)参考文献 HAO, Weituo et al. , WAFFLe: Weight Anonymized Factorization for Federated Learning , a
arXiv.org, [online] , 2020年08月13日 , [検索日 2024.08.22], Retrieved from the Internet:
URL: <https://arxiv.org/pdf/2008.05687.pdf>
YANG, Brandon et al. , CondConv: Conditionally Parameterized Convolutions for Efficient In
ference , arXiv.org, [online] , 2020年09月04日 , [検索日 2024.08.22], Retrieved from the I
nternet: URL: <https://arxiv.org/pdf/1904.04971.pdf>
CHENG, Xin et al. , Real-time Human Activity Recognition Using Conditionally Parametrized
Convolutions on Mobile and Wearable Devices , arXiv.org, [online] , 2020年06月13日 , [
検索日 2024.08.22], Retrieved from the Internet: URL: <https://arxiv.org/pdf/2006.03259.pdf>
(58)調査した分野 (Int.Cl. , D B 名)
G 0 6 N 2 0 / 0 0
G 0 6 N 3 / 0 9 8