

(12) 特許協力条約に基づいて公開された国際出願

(19) 世界知的所有権機関
国際事務局

(43) 国際公開日
2020年1月2日(02.01.2020)



(10) 国際公開番号

WO 2020/003849 A1

(51) 国際特許分類:

G06N 3/08 (2006.01)

(21) 国際出願番号:

PCT/JP2019/020906

(22) 国際出願日:

2019年5月27日(27.05.2019)

(25) 国際出願の言語:

日本語

(26) 国際公開の言語:

日本語

(30) 優先権データ:

特願 2018-119727 2018年6月25日(25.06.2018) JP

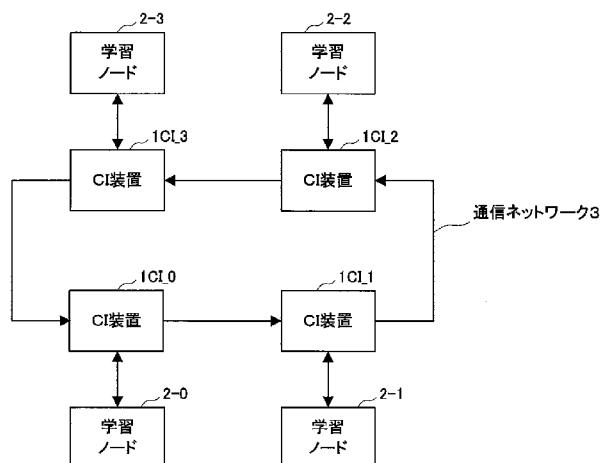
(71) 出願人: 日本電信電話株式会社 (NIPPON TELEGRAPH AND TELEPHONE CORPORATION) [JP/JP]; 〒1008116 東京都千代田区大手町一丁目5番1号 Tokyo (JP).

(72) 発明者: 加藤 順一 (KATO, Junichi); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP). 川合 健治 (KAWAI, Kenji); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP). ゴーフイクー (NGO, Huycu); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP). 有川 勇輝 (ARIKAWA, Yuki); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP). 伊藤 猛 (ITO, Tsuyoshi); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP). 坂本 健 (SAKAMOTO, Takeshi); 〒1808585 東京都武蔵野市緑町3丁目9-11 NTT 知的財産センタ内 Tokyo (JP).

(54) Title: DISTRIBUTED DEEP LEARNING SYSTEM, DISTRIBUTED DEEP LEARNING METHOD, AND COMPUTING INTERCONNECT DEVICE

(54) 発明の名称: 分散深層学習システム、分散深層学習方法、およびコンピューティングインタコネクト装置

[図1]



3... COMMUNICATION NETWORK
1CI_0, 1CI_1, 1CI_2, 1CI_3... CI DEVICE
2-0, 2-1, 2-2, 2-3... LEARNING NODE

(57) Abstract: The purpose of the present invention is to provide a distributed deep learning system that can perform, at a higher speed, cooperative processing between learning nodes connected by a communication network, while increasing speed by parallel processing of learning using a large number of learning nodes connected by the communication network. This distributed deep learning system is provided with: a plurality of computing interconnect devices 1 mutually connected via a ring-type communication network 3 in which communication is possible in one direction; and a plurality of learning nodes 2 connected in a one-to-one relationship with each of the plurality of computing interconnect devices 1. The computing interconnect devices 1 simultaneously execute, in parallel, transmission/reception processing of communication packets between the learning nodes 2 and all-reduce processing.

(74) 代理人: 山川 茂樹, 外(YAMAKAWA, Shigeki et al.); 〒1006104 東京都千代田区永田町 2 丁目 1 1 番 1 号 山王パークタワー 4 階 山川国際特許事務所内 Tokyo (JP).

(81) 指定国(表示のない限り、全ての種類の国内保護が可能): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(84) 指定国(表示のない限り、全ての種類の広域保護が可能): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), ユーラシア (AM, AZ, BY, KG, KZ, RU, TJ, TM), ヨーロッパ (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

添付公開書類:

一 国際調査報告 (条約第21条(3))

(57) 要約: 通信ネットワークに接続された多数の学習ノードによって学習を並列処理して高速化を図りつつ、通信ネットワークで接続された各学習ノード間での協調処理をより高速に行うことができる分散深層学習システムを提供することを目的とする。分散深層学習システムは、1 方向に通信可能なリング型の通信ネットワーク 3 を介して互いに接続された複数のコンピューティングインタコネクト装置 1 と、複数のコンピューティングインタコネクト装置 1 のそれぞれと一対一に接続された複数の学習ノード 2 とを備え、コンピューティングインタコネクト装置 1 は、各学習ノード 2 間の通信パケットの送受信処理と A l l - r e d u c e 処理とを同時並行して実行する。

明 細 書

発明の名称：

分散深層学習システム、分散深層学習方法、およびコンピューティングインタコネクタ装置

技術分野

[0001] 本発明は、ニューラルネットワークを用いた機械学習である深層学習を複数の学習ノードで分散協調して実行する分散学習システム、分散学習方法、およびコンピューティングインタコネクタ装置に関する。

背景技術

[0002] 様々な情報、データに対する機械学習の活用により、サービスの高度化・付加価値の提供が盛んに行われている。その際の機械学習には大きな計算リソースが必要である場合が多い。特に、深層学習と呼ばれるニューラルネットワークを用いた機械学習においては、ニューラルネットワークの構成パラメータを最適化する工程である学習において、大量の学習用データを処理する必要がある。この学習処理を高速化するために、複数の演算装置で並列処理することが1つの解決法になる。

[0003] 例えば、非特許文献1には、図26のように、4台の学習ノード300-1～300-4と、インフィニバンドスイッチ301と、ヘッドノード302とがインフィニバンドネットワーク（InfiniBand network）を介して接続された分散深層学習システムが開示されている。各学習ノード300-1～300-4には、それぞれ4台のGPU（Graphics Processing Unit）が搭載されている。この非特許文献1に開示された分散深層学習システムでは、4台の学習ノード300-1～300-4によって、学習演算を並列処理することによって高速化を図っている。

[0004] 非特許文献2には、8台のGPUを搭載した学習ノード（GPUサーバ）とイーサネット（登録商標）スイッチとがイーサネットネットワークを介し

て接続された構成が開示されている。この非特許文献2には、学習ノードを1台、2台、4台、8台、16台、32台、44台用いた場合の例がそれぞれ開示されている。非特許文献2に開示されたシステム上で、分散同期確率的勾配降下法（Distributed synchronous SGD（Stochastic Gradient Descent））を用いて機械学習を行う。具体的には、以下の手順で行う。

[0005] (I) 学習データの一部を抜き出す。抜き出した学習データの集合をミニバッチと呼ぶ。

(II) ミニバッチをGPUの台数分に分けて、各GPUに割り当てる。

(III) 各GPUにおいて、(II)で割り当てられた学習データを入力した場合のニューラルネットワークからの出力値が、正解（教師データと呼ぶ）からどれだけ乖離しているかの指標となる損失関数 $L(w)$ を求める。この損失関数を求める工程では、ニューラルネットワークの入力側の層から出力側の層に向かって順番に出力値を計算していくことから、この工程を順伝搬（forward propagation）と呼ぶ。

[0006] (IV) 各GPUにおいて、(III)で求めた損失関数値に対するニューラルネットワークの各構成パラメータ（ニューラルネットワークの重み等）による偏微分値（勾配）を求める。この工程では、ニューラルネットワークの出力側の層から入力側の層に向かって順番に各層の構成パラメータに対する勾配を計算していくことから、この工程を逆伝搬（back propagation）と呼ぶ。

(V) 各GPU毎に計算した勾配の平均を計算する。

[0007] (VI) 各GPUにおいて、(V)で計算した勾配の平均値を用いて、確率的勾配降下法（SGD：Stochastic Gradient Descent）を用いて、損失関数 $L(w)$ がより小さくなるように、ニューラルネットワークの各構成パラメータを更新する。確率的勾配降下法は、各構成パラメータの値を勾配の方向に微少量変更することにより、損失関数 $L(w)$ を小さくするという計算処理である。この処理を繰り返すことによっ

て、ニューラルネットワークは、損失関数 $L(w)$ が小さい、すなわち、正解に近い出力をする精度の高いものに更新されていく。

[0008] また、非特許文献3には、8台のGPUを搭載した学習ノード128台がインフィニバンドネットワーク(InfiniBand network)を介して接続された構成の分散深層学習システムが開示されている。

[0009] 非特許文献1～3のいずれの分散深層学習システムにおいても、学習ノード数が増えるに従い、学習速度が上がり、学習時間を短縮できることが示されている。この場合、各学習ノードで算出した勾配等のニューラルネットワーク構成パラメータの平均値を計算するため、これらの構成パラメータを学習ノード間で送受信するか、あるいは学習ノードと非特許文献1のヘッドノードとの間で送受信することにより、平均値算出等の計算を行う必要がある。

[0010] 一方、並列処理数を増やすために、ノード数を増やすにつれ、必要な通信処理は急速に増大する。従来技術のように、学習ノードやヘッドノード上で平均値算出等の演算処理やデータの送受信処理をソフトウェアで行う場合、通信処理に伴うオーバーヘッドが大きくなり、学習効率を十分に上げることが難しくなるという課題があった。

[0011] 非特許文献3には、学習処理を100サイクル行うのにかかる所要時間とこのうちの通信にかかる時間と、GPU数との関係が開示されている。この関係によると、GPU数が増えるにつれて通信にかかる時間が増えており、特にGPU数が512以上のところで急激に増加している。

先行技術文献

非特許文献

[0012] 非特許文献1: Rengan Xu and Nishanth Dandapanthu., “NVIDIA (登録商標) Tesla (登録商標) P100 GPUによるディープラーニングのパフォーマンス”, デル株式会社, 2016年, インターネット<<http://ja.community.dell.com/techcenter/m/mediagall>

ery/3765/download>

非特許文献2: Priya Goyal, Piotr Dollár, Ross Girshick, Pieter Noordhuis, Lukasz Wesolowski, Aapo Kyrola, Andrew Tulloch, Yangqing Jia, Kaiming He, “Accurate, Large Minibatch SGD: Training ImageNet in 1 Hour”, 米国コーネル大学ライブラリー, arXiv:1706.02677, 2017, インターネット<<https://arxiv.org/abs/1706.02677>>

非特許文献3: Takuya Akiba, Shuji Suzuki, Keisuke Fukuda, “Extremely Large Minibatch SGD: Training ResNet-50 on ImageNet in 15 Minutes”, 米国コーネル大学ライブラリー, arXiv:1711.04325, 2017, インターネット<<https://arxiv.org/abs/1711.04325>>

発明の概要

発明が解決しようとする課題

- [0013] 本発明は、上述した課題を解決するためになされたものであり、通信ネットワークに接続された多数の学習ノードによって学習を並列処理して高速化を図りつつ、通信ネットワークで接続された各学習ノード間での協調処理をより高速に行うことができる分散深層学習システム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

- [0014] 上述した課題を解決するために、本発明に係る分散深層学習システムは、1方向に通信可能なリング型の通信ネットワークを介して互いに接続された複数のコンピューティングインタコネクト装置と、前記複数のコンピューティングインタコネクト装置のそれぞれと一対一に接続された複数の学習ノードとを備え、各コンピューティングインタコネクト装置は、自装置に接続さ

れた学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信部と、自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信部と、この第2受信部が受信した前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信部によって取得された前記転送データを振り分ける第1振分部と、前記第1受信部が受信した前記パケットに含まれる前記受信完了フラグと前記役割とに応じて、前記第1受信部によって取得された前記ノードデータを振り分ける第2振分部と、前記第2振分部によって振り分けられた前記ノードデータ、または前記第1振分部により振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信する第1送信部と、前記第1振分部によって振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信部とを備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記第1送信部および前記第2送信部に振り分け、前記受信完了フラグがパケットの受信の完了を示し、かつ、前記役割が親である場合には、前記転送データを廃棄し、前記第2振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記ノードデータを前記第1送信部に振り分け、各学習ノードは、学習データの入力に対して演算結果を出力するニューラルネットワークと、データをパケット化して、自ノードと接続されたコンピューティングインタコネクト装置に送信する第3送信部と、自ノードと接続された前記コンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された前記転送データを取得する第3受信部と、この第3受信部が取得した前記転送データに基づいて前記ニューラルネットワークの構成パラメータデータを更新する構成パラメータ更新部と

を備えることを特徴とする。

[0015] また、本発明に係る分散深層学習システムにおいて、前記コンピューティングインタコネクト装置は、前記第1振分部によって振り分けられた前記転送データと、前記第2振分部によって振り分けられた前記ノードデータとを入力とする演算を行う演算器をさらに備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記転送データを前記演算器に振り分け、前記第2振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記ノードデータを前記演算器に振り分け、前記演算器は、前記転送データおよび前記ノードデータを入力とする演算結果を前記第1送信部に出力してもよい。

[0016] また、本発明に係る分散深層学習システムにおいて、前記コンピューティングインタコネクト装置は、前記ノードデータを記憶する機能をもつ構成パラメータメモリと、前記第1振分部によって振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として、更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されたデータを更新する構成パラメータ更新演算部と、をさらに備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記構成パラメータ更新演算部に振り分け、前記構成パラメータ更新演算部は、算出した前記更新後の構成パラメータデータを前記第1送信部および前記第2送信部に出力し、前記第1送信部は、前記更新後の構成パラメータデータをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信し、前記第2送信部は、前記更新後の構成パラメータデータをパケット化して、自装置と接続された前記学習ノードに送信してもよい。

[0017] また、本発明に係る分散深層学習方法は、1方向に通信可能なリング型の通信ネットワークを介して互いに接続された複数のコンピューティングインタコネクト装置と、前記複数のコンピューティングインタコネクト装置のそ

れぞれと一対一に接続された複数の学習ノードとを備える分散深層学習システムにおける分散深層学習方法であって、各コンピューティングインタコネクト装置が、自装置に接続された学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信ステップと、自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信ステップと、この第2受信ステップで受信された前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信ステップで取得された前記転送データを振り分ける第1振分ステップと、前記第1受信ステップで受信された前記パケットに含まれる前記受信完了フラグと前記役割とに応じて、前記第1受信ステップで取得された前記ノードデータを振り分ける第2振分ステップと、前記第2振分ステップで振り分けられた前記ノードデータ、または前記第1振分ステップで振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信する第1送信ステップと、前記第1振分ステップで振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信ステップとを備え、各学習ノードが、ニューラルネットワークに学習データを入力して演算結果を出力するニューラルネットワーク演算ステップと、データをパケット化して、自ノードと接続された前記コンピューティングインタコネクト装置に送信する第3送信ステップと、自ノードと接続された前記コンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された前記転送データを取得する第3受信ステップと、この第3受信ステップで取得された前記転送データに基づいて前記ニューラルネットワークの構成パラメータデータを更新する構成パラメータ更新ステップとを備えることを特徴とする。

[0018] また、本発明に係る分散深層学習方法において、前記コンピューティング

インタコネクト装置が、前記ノードデータを構成パラメータメモリに記憶する構成パラメータ記憶ステップと、前記第1振分ステップで振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として、更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されたデータを更新する構成パラメータ更新演算ステップと、をさらに備えていてもよい。

[0019] また、本発明に係るコンピューティングインタコネクト装置は、1方向に通信可能なリング型の通信ネットワークを介して互いに接続され、かつ、複数の学習ノードのそれぞれと一対一に接続された複数のコンピューティングインタコネクト装置であって、自装置に接続された学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信部と、自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信部と、この第2受信部が受信した前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信部によって取得された前記転送データを振り分ける第1振分部と、前記第1受信部が受信した前記パケットに含まれる前記受信完了フラグと前記役割とに応じて、前記第1受信部によって取得された前記ノードデータを振り分ける第2振分部と、前記第2振分部によって振り分けられた前記ノードデータ、または前記第1振分部により振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信する第1送信部と、前記第1振分部によって振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信部とを備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記第1送信部および前記第2送信部に振り分け、前記受信完了フラグがパケットの受信の完了を示し、かつ、前記役割が親で

ある場合には、前記転送データを廃棄し、前記第2振分部は、記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記ノードデータを前記第1送信部に振り分けることを特徴とする。

[0020] また、本発明に係るコンピューティングインタコネクタ装置において、前記第1振分部によって振り分けられた前記転送データと、前記第2振分部によって振り分けられた前記ノードデータとを入力とする演算を行う演算器をさらに備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記転送データを前記演算器に振り分け、前記第2振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記ノードデータを前記演算器に振り分け、前記演算器は、前記転送データおよび前記ノードデータを入力とする演算結果を前記第1送信部に出力してもよい。

[0021] また、本発明に係るコンピューティングインタコネクタ装置において、前記ノードデータを記憶する機能をもつ構成パラメータメモリと、前記第1振分部によって振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されているデータを更新する構成パラメータ更新演算部とをさらに備え、前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記構成パラメータ更新演算部に振り分け、前記構成パラメータ更新演算部は、算出した前記更新後の構成パラメータデータを前記第1送信部および前記第2送信部に出力し、前記第1送信部は、前記更新後の構成パラメータデータをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクタ装置に送信し、前記第2送信部は、前記更新後の構成パラメータデータをパケット化して、自装置と接続された前記学習ノードに送信してもよい。

発明の効果

[0022] 本発明によれば、コンピューティングインタコネクタ装置が各学習ノード

で計算された勾配の値の和を計算し、計算した結果を各学習ノードに返送する処理を行い、この処理と各学習ノード間の通信パケットの送受信処理とが同時並行して実行される。そのため、通信ネットワークに接続された多数の学習ノードによって学習を並列処理して高速化を図りつつ、通信ネットワークで接続された各学習ノード間での協調処理をより高速に行うことができる。

図面の簡単な説明

[0023] [図1]図1は、本発明の第1の実施の形態に係る分散深層学習システムの構成を示すブロック図である。

[図2]図2は、2層ニューラルネットワークの構成を示すブロック図である。

[図3]図3は、従来の分散深層学習処理の手順を説明する図である。

[図4]図4は、本発明の第1の実施の形態に係る学習ノードの構成を示すブロック図である。

[図5]図5は、本発明の第1の実施の形態に係る分散学習処理の手順を説明する図である。

[図6]図6は、本発明の第1の実施の形態に係る分散学習処理の手順を説明する図である。

[図7A]図7Aは、本発明の第1の実施の形態に係る分散深層学習処理システムの動作を説明する図である。

[図7B]図7Bは、本発明の第1の実施の形態に係る分散深層学習処理システムの動作を説明する図である。

[図7C]図7Cは、本発明の第1の実施の形態に係る分散深層学習処理システムの動作を説明する図である。

[図7D]図7Dは、図7Aから図7Cの演算情報テーブルを示す図である。

[図8A]図8Aは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8B]図8Bは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8C]図8Cは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8D]図8Dは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8E]図8Eは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8F]図8Fは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図8G]図8Gは、本発明の第1の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図9]図9は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置の構成を示すブロック図である。

[図10]図10は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0）の動作を説明する図である。

[図11]図11は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__1）の動作を説明する図である。

[図12]図12は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__2）の動作を説明する図である。

[図13]図13は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0）の動作を説明する図である。

[図14]図14は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__1）の動作を説明する図である。

[図15]図15は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__2）の動作を説明する図である。

[図16]図16は、本発明の第1の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0）の動作を説明する図である。

[図17A]図17Aは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17B]図17Bは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17C]図17Cは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17D]図17Dは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17E]図17Eは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17F]図17Fは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図17G]図17Gは、本発明の第2の実施の形態に係る分散深層学習システムの動作を説明する図である。

[図18]図18は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置の構成を示すブロック図である。

[図19]図19は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0'）の動作を説明する図である。

[図20]図20は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__1'）の動作を説明する図である。

[図21]図21は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__2'）の動作を説明する図である。

[図22]図22は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0'）の動作を説明する図である。

[図23]図23は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__1'）の動作を説明する図である。

[図24]図24は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__2'）の動作を説明する図である。

[図25]図25は、本発明の第2の実施の形態に係るコンピューティングインタコネクタ装置（1C1__0'）の動作を説明する図である。

[図26]図26は、従来の分散深層学習システムの構成を示すブロック図である。

発明を実施するための形態

[0024] 以下、本発明の好適な実施の形態について、図1から図25を参照して詳細に説明する。

[0025] [第1の実施の形態]

図1は、本発明の第1の実施の形態に係る分散深層学習システムの構成を示すブロック図である。本実施の形態に係る分散深層学習システムは、複数のコンピューティングインタコネクト (Computing Interconnect: CI) 装置1CI__0~1CI__3がリング型の通信ネットワーク3で互いに接続され、コンピューティングインタコネクト装置1CI__0~1CI__3のそれぞれに学習ノード2-0~2-3のそれぞれが一对一となるように接続された構造を有する。

[0026] なお、本発明において、コンピューティングインタコネクト装置とは、ネットワーク上に分散配置されている機器を意味する。以下、コンピューティングインタコネクト装置1CI__0~1CI__3を総称してコンピューティングインタコネクト1ということがある。同様に、学習ノード2-0~2-3を総称して学習ノード2ということがある。

[0027] 学習ノード2は、例えば、CPU (Central Processing Unit)、GPU (Graphics Processing Unit) 等の演算資源、記憶装置およびインタフェースを備えたコンピュータと、これらのハードウェア資源を制御するプログラムによって実現してもよいし、FPGA (Field Programmable Gate Array) やASIC (Application Specific Integrated Circuit) に形成したLSI (Large Scale Integration) 回路で実現してもよい。

[0028] コンピューティングインタコネクト装置1CI__0~1CI__3間は、イーサネット (登録商標) や、インフィニバンド (InfiniBand、登

録商標)などの、通信パケットをやりとりすることで通信を行う通信ネットワークで互いに接続されている。

[0029] コンピューティングインタコネクト装置1C1__0~1C1__3と学習ノード2-0~2-3との間は、イーサネットや、インフィニバンド(InfiniBand)などの通信ネットワークで互いに接続されていてもよい。あるいは、学習ノード2-0~2-3内のPCI Express(登録商標)などのI/Oインタフェースにコンピューティングインタコネクト装置1C1__0~1C1__3を直に挿入する接続構成を採用してもよい。

[0030] [学習ノードの説明]

学習ノード2は、数学モデルであるニューラルネットワークの出力値を計算し、さらに、学習データに応じてニューラルネットワークの構成パラメータを更新して出力値の精度を向上させていく学習機能をもつ装置である。ニューラルネットワークは、各学習ノード2-0~2-3内に構築される。なお、学習ノード2-0~2-3が備える各機能ブロックの詳細については後述する。

[0031] [学習についての説明]

学習ノード2におけるニューラルネットワークの学習処理について、教師データ付き学習を例に説明する。図2にニューラルネットワークの例として入力層(第1層)、中間層(第2層)、出力層(第3層)からなるごく単純な2層ニューラルネットワークを示す。図2の $N_k(i)$ は第 k 層、 i 番目のニューロンである。 x_1, x_2 は入力、 y_1, y_2 は出力、 $w_1(11), w_1(12), \dots, w_1(23)$ は第1層目の重みパラメータ、 $w_2(11), w_2(12), \dots, w_2(32)$ は第2層目の重みパラメータである。

[0032] 教師データ付き学習の場合、各学習データには対応する教師データ(正解データ)が予め用意されており、ニューラルネットワークの出力値が教師データに近くなるように、ニューラルネットワークの構成パラメータを更新していく。図2の例の場合のニューラルネットワークの構成パラメータは、重

み $w_1(11)$, $w_1(12)$, $\dots$, $w_1(23)$, $w_2(11)$, $w_2(12)$, $\dots$, $w_2(32)$ である。これらの構成パラメータを最適化していくことにより、ニューラルネットワークの精度を上げていく。

[0033] 具体的には、ニューラルネットワークの出力値が教師データとどれだけ乖離しているかの指標となる損失関数を定め、この損失関数が小さくなるように構成パラメータを更新していく。この例では、入学習データ x_1 , x_2 に対応する出力値を y_1 , y_2 、教師データを t_1 , t_2 とすると、損失関数 L は、例えば次式のようにになる。

[0034] [数1]

$$L = \frac{1}{2} \sum_{k=1}^2 (y_k - t_k)^2 \quad \dots (1)$$

[0035] 次に、この損失関数 L に対するニューラルネットワークの各構成パラメータによる偏微分値を成分とするベクトル（これを勾配と呼ぶ）を求める。この例では、勾配は以下のようにになる。

[0036] [数2]

$$\left(\frac{\partial L}{\partial w_1(11)}, \frac{\partial L}{\partial w_1(12)}, \dots, \frac{\partial L}{\partial w_1(23)}, \frac{\partial L}{\partial w_2(11)}, \frac{\partial L}{\partial w_2(12)}, \dots, \frac{\partial L}{\partial w_2(32)} \right) \quad \dots (2)$$

[0037] 次に、勾配を用いて、損失関数 L がより小さくなるように、ニューラルネットワークの各構成パラメータを更新する。更新の方法はいろいろあるが、例えば勾配降下法を用いて、それぞれの重みパラメータを以下のように更新する。

[0038] [数3]

$$\begin{aligned} w_1(11) &\leftarrow w_1(11) - \eta \frac{\partial L}{\partial w_1(11)} \\ &\dots \dots \dots \\ w_2(32) &\leftarrow w_2(32) - \eta \frac{\partial L}{\partial w_2(32)} \quad \dots (3) \end{aligned}$$

[0039] ここで、 η は学習率と呼ばれる定数である。式 (3) により、各重みパラメータを、勾配と逆の方向、すなわち、損失関数 L を減少させる方向に学習

率 η に比例する量だけ変化させている。そのため、更新後のニューラルネットワークの損失関数 L は更新前より小さくなる。

[0040] このように、1組の入力学習データに対して、損失関数 L の計算、勾配の計算、構成パラメータの更新の処理を行う。そして、この構成パラメータの更新されたニューラルネットワークに対して、次の入力学習データを入力して同じ処理を行い、構成パラメータを更新する。このサイクルを繰り返すことにより、損失関数 L が小さいニューラルネットワークに更新していくことで、ニューラルネットワークの学習を行う。

[0041] ここで、損失関数 L を求める工程では、ニューラルネットワークの入力層から出力層に向かって順番に出力値を計算していくことから、この工程を順伝搬（forward propagation）と呼ぶ。一方、勾配を求める工程では、ニューラルネットワークの出力層から入力層に向かって順番に各層の構成パラメータに対する勾配を計算していく逆伝搬（back propagation）と呼ぶ手法を用いることが多い。

[0042] [複数の学習ノードによる分散学習処理]

以上のようなニューラルネットワークの学習で十分な精度を達成するには、大量の学習データをニューラルネットワークに入力して学習処理を繰り返す必要があり、長い時間を要する。この学習にかかる所要時間を短縮することは大きなメリットがある。

[0043] 学習にかかる所要時間を短縮するため、同じニューラルネットワークの学習ノード2を複数用意して、学習データをそれぞれの学習ノード2に分けて並列で学習させることにより、トータルの学習時間を短縮する分散協調学習の手法がとられる。

[0044] 以下、従来の分散学習処理の手順を図3を用いて説明する。

最初に、学習データ x を学習ノード400-0～400-3の台数分に分けて、各学習ノード400-0～400-3に割り当てる。なお、図3では、各学習ノード400-0～400-3に割り当てる学習データの代表として $x_0 \sim x_3$ を1つずつ記載しているが、学習データ $x_0 \sim x_3$ はそれぞれ

1または複数の学習データの集合からなる。

[0045] 次に、各学習ノード400-0～400-3は、それぞれ学習データ $x_0 \sim x_3$ をニューラルネットワークに入力して順伝搬（forward propagation）の手法によりそれぞれ損失関数 L を求める（図3ステップS100）。なお、得られる損失関数 L は、各学習ノード400-0～400-3（各ニューラルネットワーク）につき1つである。

[0046] 続いて、各学習ノード400-0～400-3は、ステップS100で求めた損失関数 L の勾配を逆伝搬（back propagation）の手法により求める（図3ステップS101）。損失関数 L の勾配とは、式（2）に示すように構成パラメータ毎の成分を含むベクトルである。

[0047] 次に、各学習ノード400-0～400-3でそれぞれ計算した勾配の平均を例えばヘッドノード402において計算して、計算した結果をヘッドノード402から各学習ノード400-0～400-3に返送する（図3ステップS102）。この処理を「All-reduce処理」と呼ぶ。

なお、勾配の平均の代わりに勾配の和を計算するようにしてもよい。このとき、例えば、次の重みパラメータの更新処理時の学習率 η に（ $1/\text{学習ノード数}$ ）を乗じれば、勾配の平均値を求めるのと同じ結果になる。

さらに、勾配の平均の代わりに、各勾配に重みづけ定数をかけて重み付き平均を用いるようにしてもよいし、各勾配の二乗の和をとるようにしてもよい。

[0048] 最後に、各学習ノード400-0～400-3は、ステップS102で計算された勾配の平均値を用いて、ニューラルネットワークの重みパラメータを更新する（図3ステップS103）。

以上で、分散学習の1サイクルが終了する。

[0049] [学習ノードの機能ブロック]

次に、本実施の形態に係る分散深層学習システムの動作の概要についての説明に先立って学習ノード2の機能構成を説明する。図4は学習ノード2の構成例を示すブロック図である。

学習ノード2は、入力部20、損失関数計算部21、勾配計算部22、送信部23、受信部24、構成パラメータ更新部25、およびニューラルネットワーク26を備える。なお、学習ノード2-0~2-3はそれぞれ同様の構成を有する。

[0050] 入力部20は、学習データを受け取る。

損失関数計算部21は、学習データが入力されたときに、損失関数 L をニューラルネットワーク26の構成パラメータ毎および学習データ毎に計算する。

勾配計算部22は、損失関数 L の勾配を学習データ毎に計算した後に、勾配を集計した値を構成パラメータ毎に生成する。

[0051] 送信部23（第3送信部）は、勾配計算部22によって計算された勾配の値をパケット化してコンピューティングインタコネクタ装置1に送信する。より詳細には、送信部23は、勾配計算部22によって計算された勾配の計算結果と、この計算結果に対応する構成パラメータに固有のシーケンシャル番号と対応する演算IDとを後述する通信パケットのデータペイロードに書き込んで、自ノードと接続されているコンピューティングインタコネクタ装置1に送信する。

[0052] 受信部24（第3受信部）は、コンピューティングインタコネクタ装置1から送信された通信パケットを受信する。より詳細には、学習ノードの受信部24は、自ノードと接続されているコンピューティングインタコネクタ装置1から受信した通信パケットのデータペイロードから勾配の和（転送データ）の計算結果とシーケンシャル番号と演算IDとを取り出す。

[0053] 構成パラメータ更新部25は、コンピューティングインタコネクタ装置1から送信された通信パケットに格納されている勾配の和を用いてニューラルネットワークの構成パラメータ（重みパラメータ）を更新する。より詳細には、構成パラメータ更新部25は、勾配の和の計算結果を基に、シーケンシャル番号で特定される、ニューラルネットワーク26の構成パラメータを更新する。

[0054] ニューラルネットワーク 26 は、数学モデルであるニューラルネットワークの出力値を計算する。本実施の形態では、1つの演算 ID の演算対象となっている各学習ノード 2 のニューラルネットワーク 26 の構成は同一であるものとし、以下の他の実施の形態でも同様とする。

[0055] [本実施の形態の分散処理]

次に、本実施の形態に係る学習ノード 2-0~2-3 とコンピューティングインタコネクタ装置 1C1__0~1C1__3 とで行われる分散学習処理の手順を図 5 を用いて説明する。本実施の形態では、各学習ノード 2-0~2-3 は、従来例と同様に、それぞれ学習データ $x_0 \sim x_3$ をニューラルネットワーク 26 に入力して損失関数 L をそれぞれ計算する（図 5 ステップ S 200）。

[0056] より詳細には、入力部 20 に学習データが入力される。その後、損失関数計算部 21 は、学習データが入力されると、損失関数 L をニューラルネットワーク 26 の構成パラメータ毎および学習データ毎に計算する。

[0057] 続いて、勾配計算部 22 は、算出された損失関数 L の勾配を計算する（図 5 ステップ S 201）。より詳細には、勾配計算部 22 は、損失関数 L の勾配を学習データ毎に計算した後に、勾配を集計した値を構成パラメータ毎に生成する。

[0058] そして、各学習ノード 2-0~2-3 の送信部 23 は、それぞれ算出された勾配の値を、各学習ノード 2-0~2-3 とそれぞれ接続されたコンピューティングインタコネクタ装置 1C1__0~1C1__3 に送信する（図 5 ステップ S 202）。

[0059] なお、図 3 の従来例と同様に、図 5 では、各学習ノード 2-0~2-3 に割り当てる学習データの代表として $x_0 \sim x_3$ を 1 つずつ記載しているが、学習データ $x_0 \sim x_3$ はそれぞれ 1 または複数の学習データの集合からなる。

[0060] 次に、コンピューティングインタコネクタ装置 1C1__0~1C1__3 は、各学習ノード 2-0~2-3 の送信部 23 から送信された各勾配の計算値

を、コンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 間の通信ネットワーク 3 を使って順番に加算していく。この結果得られる全勾配の平均値を各学習ノード 2 - 0 ~ 2 - 3 に送信する（図 5 ステップ S 2 0 3, S 2 0 4）。このように本実施の形態では、コンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 を用いて *All-reduce* 処理を行う。

[0061] 最後に、各学習ノード 2 - 0 ~ 2 - 3 の受信部 2 4 は、コンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 から送信された勾配の平均値を受信する。構成パラメータ更新部 2 5 は、受信された勾配の平均値を用いて、ニューラルネットワーク 2 6 の構成パラメータを更新する（図 5 ステップ S 2 0 5）。

[0062] なお、勾配の平均の代わりに勾配の和を計算するようにしてもよい。このとき、例えば、次の重みパラメータの更新処理時の学習率 η に（1 / 学習ノード数）を乗じれば、勾配の平均値を求めるのと同じ結果になる。また、各勾配に重みづけ定数をかけて重み付き平均を用いるようにしてもよいし、勾配の二乗平均平方根をとるようにしてもよい。

以上で、本実施の形態の分散学習の 1 サイクルが終了する。

[0063] 通常、勾配計算は逆伝搬の手法に従って、ニューラルネットワーク 2 6 の出力層から入力層に向かって順番に各層の構成パラメータ（重みパラメータ）に対する勾配の成分を計算していく。したがって、各学習ノード 2 - 0 ~ 2 - 3 の勾配計算結果をコンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 に送信するにあたっては、全ての層の勾配計算が終わるまで待つ必要はない。

[0064] そこで、図 6 に示すように、各学習ノード 2 - 0 ~ 2 - 3 の損失関数計算部 2 1 は、まず上記と同様に損失関数 L を計算し（図 6 ステップ S 2 0 0）、損失関数 L の勾配を計算する（図 6 ステップ S 2 0 1）。その後、ステップ S 2 0 1 において勾配計算部 2 2 がすべての構成パラメータに対する勾配成分の計算が終了するのを待つことなく、送信部 2 3 は計算が終わった構成

パラメータに対する勾配成分からコンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 に送信することができる（図 6 ステップ S 2 0 6）。

[0065] コンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 は、各学習ノード 2 - 0 ~ 2 - 3 から送信された勾配成分の平均値を計算し（図 6 ステップ S 2 0 7）、計算が終わった勾配成分の平均値を各学習ノード 2 - 0 ~ 2 - 3 に送信する（図 6 ステップ S 2 0 8）。

[0066] 各学習ノード 2 - 0 ~ 2 - 3 の受信部 2 4 は、コンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 から計算結果を受信すると、構成パラメータ更新部 2 5 は、すべての計算結果を受信するまで待つことなく、受信した勾配成分の平均値を用いて、対応する構成パラメータを更新する（図 6 ステップ S 2 0 9）。

こうして、本実施の形態では勾配計算と `All-reduce` 処理と構成パラメータ更新とをパイプライン式に処理できるので、学習処理の更なる高速化が可能である。

[0067] [分散深層学習システムの動作の概要]

図 7 A、図 7 B、図 7 C は、本実施の形態の分散深層学習システムの典型的な動作例を説明する図である。

分散深層学習システムでは、ニューラルネットワークモデル、入力学習データがそれぞれ異なる 4 つの学習演算（演算 1 ~ 4）を順番に処理している。

[0068] 図 7 A に示すように、まず、演算 1 が学習ノード 2 - 0 ~ 2 - 3 の 4 台の学習ノード 2 - 0 ~ 2 - 3 で並列演算される。これを終了後、図 7 B に示すように、演算 2 が学習ノード 2 - 0、学習ノード 2 - 1、学習ノード 2 - 3 の 3 台の学習ノードで並列演算される。

[0069] 最後に、図 7 C に示すように演算 3 と演算 4 が同時に処理されている。このとき、演算 3 は学習ノード 2 - 0 および学習ノード 2 - 1 で、演算 4 は学習ノード 2 - 2 および学習ノード 2 - 3 で並列処理されている。

[0070] 本発明の分散深層学習システムでは、並列学習演算している学習ノード 2

に接続するコンピューティングインタコネクト装置 1 C I __ 0 ~ 1 C I __ 3 のうち、1 台が親として動作し、それ以外は子として動作する。図 7 A の例では、コンピューティングインタコネクト装置 1 C I __ 0 が親、コンピューティングインタコネクト装置 1 C I __ 1 ~ 1 C I __ 3 が子として動作している。

[0071] また、図 7 B では、コンピューティングインタコネクト装置 1 C I __ 0 が親、コンピューティングインタコネクト装置 1 C I __ 1, 1 C I __ 3 が子として動作している。学習データの規模が大きい場合やそれほど高速に処理する必要がない場合は、図 7 B のように、リング型の通信ネットワーク 3 に接続される学習ノード 2 - 0 ~ 2 - 3 の一部だけを使うことも想定される。

[0072] 図 7 C では、演算 3 に対しては、コンピューティングインタコネクト装置 1 C I __ 0 を親、コンピューティングインタコネクト装置 1 C I __ 1 を子としている。また、演算 4 に対しては、コンピューティングインタコネクト装置 1 C I __ 2 を親、コンピューティングインタコネクト装置 1 C I __ 3 を子として動作させている。

[0073] このような動作を進めるにあたって、各演算ごとに、各コンピューティングインタコネクト装置 1 が、「親」、「子」、「対象外」のいずれの役割を割り当てられているか示す演算情報テーブル（図 7 D）を用いる。各学習ノード 2 - 0 ~ 2 - 3 およびコンピューティングインタコネクト装置 1 C I __ 1 ~ 1 C I __ 3 はこの演算情報テーブルを共有し、この内容に従って、各演算ごとに指定された動作を行う。

[0074] [分散深層学習システムの動作の具体例]

以下、図 7 B の場合（演算 1 D = 2 の場合）を例にとって、図 8 A から図 8 G を用いて本実施の形態の分散深層学習システムの動作を説明する。

[0075] 図 8 A に示すように、コンピューティングインタコネクト装置 1 C I __ 0 に接続された学習ノード 2 - 0 から勾配の計算結果 G 0 をコンピューティングインタコネクト装置 1 C I __ 0 に送信する。コンピューティングインタコ

ネクト装置 1C1__0 は、勾配の計算結果 G_0 を通信ネットワーク 3 を介してコンピューティングインタコネクト装置 1C1__1 に転送する。

[0076] 図 8B に示すように、コンピューティングインタコネクト装置 1C1__1 は、コンピューティングインタコネクト装置 1C1__0 から送信された勾配の計算結果 G_0 と、コンピューティングインタコネクト装置 1C1__1 の直下の学習ノード 2-1 から送信された勾配の計算結果 G_1 との和 $G_0 + G_1$ を計算する。コンピューティングインタコネクト装置 1C1__1 は、この計算結果 $G_0 + G_1$ を通信ネットワーク 3 を介してコンピューティングインタコネクト装置 1C1__2 に送信する。

[0077] 図 8C に示すように、本演算では、学習ノード 2-2 は演算対象外なので、コンピューティングインタコネクト装置 1C1__2 は、コンピューティングインタコネクト装置 1C1__1 から送信された勾配の計算結果 $G_0 + G_1$ に対して和演算を行わない。コンピューティングインタコネクト装置 1C1__2 は、計算結果 $G_0 + G_1$ をそのまま通信ネットワーク 3 を介してコンピューティングインタコネクト装置 1C1__3 へ送信する。

[0078] 図 8D に示すように、コンピューティングインタコネクト装置 1C1__3 では、コンピューティングインタコネクト装置 1C1__1 と同様の演算を行う。コンピューティングインタコネクト装置 1C1__3 は、コンピューティングインタコネクト装置 1C1__2 から送信された勾配の和の計算結果 $G_0 + G_1$ と、コンピューティングインタコネクト装置 1C1__3 の直下の学習ノード 2-3 から送信された勾配の計算結果 G_3 との和 $\Sigma G = G_0 + G_1 + G_3$ を計算する。コンピューティングインタコネクト装置 1C1__3 は、この計算結果 ΣG を通信ネットワーク 3 を介してコンピューティングインタコネクト装置 1C1__0 に送信する。

[0079] また、図 8D に示すように、勾配の和の計算結果 ΣG を受信したコンピューティングインタコネクト装置 1C1__0 は、受信した勾配の和 ΣG を直下の学習ノード 2-0 とコンピューティングインタコネクト装置 1C1__1 とに送信する。

[0080] 図8Eに示すように、勾配の和 ΣG を受信したコンピューティングインタコネクタ装置1C1__1は、勾配の和 ΣG を直下の学習ノード2-1とコンピューティングインタコネクタ装置1C1__2とに送信する。

[0081] 図8Fに示すように、勾配の和 ΣG を受信したコンピューティングインタコネクタ装置1C1__2は、勾配の和 ΣG を直下の学習ノード2-2へは送信せず、コンピューティングインタコネクタ装置1C1__3のみに通信ネットワーク3を介して送信する。

[0082] 図8Gに示すように、コンピューティングインタコネクタ装置1C1__3は、コンピューティングインタコネクタ装置1C1__2から送信された勾配の和 ΣG を直下の学習ノード2-3とコンピューティングインタコネクタ装置1C1__0とに送信する。

[0083] 最後に、図8Gに示すように、勾配の和 ΣG を受信したコンピューティングインタコネクタ装置1C1__0は、勾配の和 ΣG を廃棄する。

以上の動作により、各学習ノード2-0, 2-1, 2-3に勾配の和 ΣG が送信される。

[0084] [コンピューティングインタコネクタ装置の構成]

図9は、コンピューティングインタコネクタ装置1の構成例を示すブロック図である。

コンピューティングインタコネクタ装置1は、受信部100, 103、振分部101, 104、バッファメモリ102, 105、加算器106、送信部107, 108、および制御部109を備える。

[0085] 受信部100（第2受信部）は、1方向（本実施の形態では反時計回りの方向）に限定して通信を行うリング型の通信ネットワーク3において隣接する上流側のコンピューティングインタコネクタ装置1（例えば、図1においてそれぞれ左隣のコンピューティングインタコネクタ装置1）からの通信パケットを受信し、このパケットに格納されたデータ（転送データ）を取得する。

[0086] 振分部101（第1振分部）は、通信パケットに含まれるデータの受信が

完了されたか否かを示す受信完了フラグ（完了／未了）とそのコンピューティングインタコネクタ装置 1 に割り当てられた役割（親／子／計算対象外（非親子））に応じて受信部 100 からのデータを振り分ける。

バッファメモリ 102 は、振分部 101 からのデータを一時的に記憶する。

[0087] 受信部 103（第 1 受信部）は、コンピューティングインタコネクタ装置 1 の直下に設けられている学習ノード 2 からの通信パケットを受信し、このパケットに格納されたデータ（ノードデータ）を取得する。

振分部 104（第 2 振分部）は、自装置であるコンピューティングインタコネクタ装置 1 に割り当てられた役割に応じて受信部 103 からのデータを振り分ける。

[0088] バッファメモリ 105 は、振分部 104 からのデータを一時的に記憶する。

加算器 106（演算器）は、バッファメモリ 102，105 に一時的に記憶された勾配の値を読み出して勾配の和を計算する。

送信部 107（第 1 送信部）は、リング型のネットワーク 3 において隣接する下流側のコンピューティングインタコネクタ装置 1（右隣のコンピューティングインタコネクタ装置 1）へ、振分部 101 または振分部 104 によって振り分けられたデータをパケット化した通信パケットを送信する。

[0089] 送信部 108（第 2 送信部）は、コンピューティングインタコネクタ 1 の直下に設けられている学習ノード 2 に、振分部 101 によって振り分けられたデータをパケット化した通信パケットを送信する。

制御部 109 は、バッファメモリ 102，105 を制御する。

[0090] [コンピューティングインタコネクタ装置の動作の具体例]

図 10 は、図 8 A におけるコンピューティングインタコネクタ装置 1 C 1 0 の動作を説明する図である。図 10 に示すように、通信パケットは、通信ヘッダとデータペイロードとからなる。

[0091] 学習ノード 2-0 から送信される通信パケット R P 0 のデータペイロード

には、学習ノード2-0で計算された勾配値「G0」と、演算ID「002」と、勾配値のシーケンシャル番号「003」と、コンピューティングインタコネクタ装置1C1__0で勾配の和の取得完了または未了を示す受信完了フラグ（図10の例では「未了」）とが格納されている。

[0092] コンピューティングインタコネクタ装置1C1__0の受信部103は、受信した通信パケットRPOのデータペイロードから勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部104に渡す。

[0093] 振分部104では、受信部103から受け取った演算IDと演算情報テーブルをつきあわせて、この演算ID=2に対しては、コンピューティングインタコネクタ装置1C1__0は「親」として動作すべきことを識別する。これにより、振分部104は勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを送信部107に渡す。

[0094] 送信部107は、振分部104から受け取った勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC1のデータペイロードに格納する。そして、送信部107は、通信パケットTPC1を、通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__0に隣接する下流側のコンピューティングインタコネクタ装置（図8Aの例ではコンピューティングインタコネクタ装置1C1__1）へ送信する。

[0095] 図11は、図8Bにおけるコンピューティングインタコネクタ装置1C1__1の動作を示している。

コンピューティングインタコネクタ装置1C1__1の受信部100は、コンピューティングインタコネクタ装置1C1__0から受信した通信パケットTPC1のデータペイロードから勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0096] 振分部101は、受信部100から受け取った受信完了フラグが「未了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクタ装置1C1__1は「子」として動作すべきことを識別する。これにより、振分部

101は、勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとをバッファメモリ102に格納する。

[0097] 一方、コンピューティングインタコネクタ装置1C1__1の受信部103は、直下に接続されている学習ノード2-1から受信した通信パケットRPC1のデータペイロードから勾配値G1と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部104に渡す。

[0098] コンピューティングインタコネクタ装置1C1__1の振分部104では、受信部103から受け取った演算IDと演算情報テーブルをつきあわせて、この演算ID=2に対しては、コンピューティングインタコネクタ装置1C1__1は「子」として動作すべきことを識別する。これにより、振分部104は、勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとをバッファメモリ105に格納する。

[0099] コンピューティングインタコネクタ装置1C1__1の制御部109は、バッファメモリ102とバッファメモリ105に同一のシーケンシャル番号の勾配値「G0」と「G1」とが揃った時点で、バッファメモリ102から勾配値「G0」とシーケンシャル番号と受信完了フラグとを読み出す。これと同時に、制御部109は、バッファメモリ105から勾配値「G1」とシーケンシャル番号と受信完了フラグとを読み出し、勾配値「G0」と「G1」とを加算器106に渡す。

[0100] 加算器106は、勾配値「G0」と「G1」とを加算する。また、制御部109は、バッファメモリ102から読み出した演算IDとシーケンシャル番号と受信完了フラグとを送信部107に渡す。

[0101] コンピューティングインタコネクタ装置1C1__1の送信部107は、加算器106によって計算された勾配の和「G0+G1」、および制御部109から受け取った演算IDとシーケンシャル番号と受信完了フラグを通信パケットTPC2のデータペイロードに格納する。そして送信部107は、通信パケットTPC2を、通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__1に隣接する下流側のコンピューティングインタ

コネクタ装置（図 8 B の例ではコンピューティングインタコネクタ装置 1 C I __ 2）へ送信する。

- [0102] 図 1 2 は、図 8 C におけるコンピューティングインタコネクタ装置 1 C I __ 2 の動作を示している。

コンピューティングインタコネクタ装置 1 C I __ 2 の受信部 1 0 0 は、コンピューティングインタコネクタ装置 1 C I __ 1 から受信した通信パケット T P C 2 のデータペイロードから勾配値「 $G_0 + G_1$ 」と演算 I D とシーケンシャル番号と受信完了フラグとを取り出して振分部 1 0 1 に渡す。

- [0103] 振分部 1 0 1 は、受け取った演算 I D と演算情報テーブルをつきあわせてこの演算 I D = 2 に対しては、コンピューティングインタコネクタ装置 1 C I __ 2 は「演算対象外（非親子）」として動作すべきことを識別する。これにより、振分部 1 0 1 は、勾配値「 $G_0 + G_1$ 」と演算 I D とシーケンシャル番号と受信完了フラグとを送信部 1 0 7 に送信する。

- [0104] 送信部 1 0 7 は、振分部 1 0 1 から受け取った勾配値「 $G_0 + G_1$ 」と演算 I D とシーケンシャル番号と受信完了フラグとを通信パケット T P C 3 のデータペイロードに格納する。そして、送信部 1 0 7 は、通信パケット T P C 3 を通信ネットワーク 3 を介してコンピューティングインタコネクタ装置 1 C I __ 2 に隣接する下流側のコンピューティングインタコネクタ装置（図 8 C の例ではコンピューティングインタコネクタ装置 1 C I __ 3）へ送信する。

- [0105] 図 1 3 は、図 8 D におけるコンピューティングインタコネクタ装置 1 C I __ 0 の動作を示している。

コンピューティングインタコネクタ装置 1 C I __ 0 の受信部 1 0 0 は、コンピューティングインタコネクタ装置 1 C I __ 0 に隣接する上流側のコンピューティングインタコネクタ装置（図 8 D の例ではコンピューティングインタコネクタ装置 1 C I __ 3）から受信した通信パケット T P C 0 のペイロードから勾配の和 $\Sigma G (= G_0 + G_1 + G_3)$ と演算 I D とシーケンシャル番号と受信完了フラグとを取り出して振分部 1 0 1 に渡す。

- [0106] 振分部101は、受信部100から受け取った受信完了フラグが「未了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクタ装置1C1__0は「親」として動作すべきことを識別する。これにより、振分部101は、勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを送信部107、および送信部108に渡す。
- [0107] このとき、「親」であるコンピューティングインタコネクタ装置1C1__0が隣接する上流側のコンピューティングインタコネクタ装置1C1__3から受信完了フラグが「未了」である通信パケットを受信したということは、通信パケットがリング型の通信ネットワーク3を一巡し、勾配の和の計算が完了したことを意味する。そこで、コンピューティングインタコネクタ装置1C1__0の振分部101は、受信部100から受け取った受信完了フラグを、「未了」から「完了」を示す値に変更した上で送信部107と送信部108とにデータを渡す。
- [0108] コンピューティングインタコネクタ装置1C1__0の送信部107は、振分部101から受け取った勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC1のデータペイロードに格納する。そして、送信部107は、通信パケットTPC1を、通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__0に隣接する下流側のコンピューティングインタコネクタ装置（図8Dの例ではコンピューティングインタコネクタ装置1C1__1）へ送信する。
- [0109] コンピューティングインタコネクタ装置1C1__0の送信部108は、振分部101から受け取った勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTP0のデータペイロードに格納して、通信パケットTP0を学習ノード2-0へ送信する。
- [0110] 図14は、図8Eにおけるコンピューティングインタコネクタ装置1C1__1の動作を示している。
- コンピューティングインタコネクタ装置1C1__1の受信部100は、コ

ンピューティングインタコネクタ装置 1 C I __ 0 から受信した通信パケット T P C 1 のデータペイロードから勾配の和 ΣG と演算 I D とシーケンシャル番号と受信完了フラグとを取り出して振分部 1 0 1 に渡す。

[0111] 振分部 1 0 1 は、受信部 1 0 0 から受け取った受信完了フラグが「完了」を示していること、また、受け取った演算 I D と演算情報テーブルをつきあわせてこの演算 I D = 2 に対してはコンピューティングインタコネクタ装置 1 C I __ 1 は「子」として動作すべきことを識別する。これにより、振分部 1 0 1 は、勾配の和 ΣG と演算 I D とシーケンシャル番号と受信完了フラグとを送信部 1 0 7、および送信部 1 0 8 に渡す。

[0112] 送信部 1 0 7 は、振分部 1 0 4 から受け取った勾配の和 ΣG と演算 I D とシーケンシャル番号と受信完了フラグとを通信パケット T P C 2 のデータペイロードに格納する。そして送信部 1 0 7 は、通信パケット T P C 2 を、コンピューティングインタコネクタ装置 1 C I __ 1 に隣接する下流側のコンピューティングインタコネクタ装置（図 8 E の例ではコンピューティングインタコネクタ装置 1 C I __ 2）へ送信する。

[0113] コンピューティングインタコネクタ装置 1 C I __ 1 の送信部 1 0 8 は、振分部 1 0 1 から受け取った勾配の和 ΣG と演算 I D とシーケンシャル番号と受信完了フラグとを通信パケット T P 1 のデータペイロードに格納して、通信パケット T P 1 を学習ノード 2 - 1 へ送信する。

[0114] 図 1 5 は、図 8 F におけるコンピューティングインタコネクタ装置 1 C I __ 2 の動作を示している。

コンピューティングインタコネクタ装置 1 C I __ 2 の受信部 1 0 0 は、コンピューティングインタコネクタ装置 1 C I __ 1 から受信した通信パケット T P C 2 のデータペイロードから勾配の和 ΣG と演算 I D とシーケンシャル番号と受信完了フラグとを取り出して振分部 1 0 1 に渡す。

[0115] 振分部 1 0 1 は、受け取った演算 I D と演算情報テーブルをつきあわせてこの演算 I D = 2 に対しては「演算対象外（非親子）」として動作すべきことを識別する。これにより、振分部 1 0 1 は、勾配の和 ΣG と演算 I D とシ

ーケンシャル番号と受信完了フラグとを送信部107に送信する。

[0116] 送信部107は、振分部101から受け取った勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC3のデータペイロードに格納する。そして、送信部107は、通信パケットTPC3を通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__2に隣接する下流側のコンピューティングインタコネクタ装置（図8Fの例ではコンピューティングインタコネクタ装置1C1__3）へ送信する。

[0117] 図16は、図8Gにおけるコンピューティングインタコネクタ装置1C1__0の動作を示している。

コンピューティングインタコネクタ装置1C1__0の受信部100は、コンピューティングインタコネクタ装置1C1__0に隣接する上流側のコンピューティングインタコネクタ装置（図8Gの例ではコンピューティングインタコネクタ装置1C1__3）から受信した通信パケットTPC0のペイロードから勾配の和 ΣG とシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0118] 振分部101は、受信部100から受け取った受信完了フラグが「完了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクタ装置1C1__0は「親」として動作すべきことを識別する。そして、振分部101は、受信部100から受け取った勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを廃棄する。

[0119] なお、ここでは、各勾配の和を用いて重みパラメータの更新処理を行う場合を例に説明したが、各勾配の和の代わりに各勾配の重み付き和を用いる場合は、加算器106の代わりに G_{in} と G_{local} に対する重み付き和演算器を用いてもよい。また、各勾配の和の代わりに各勾配の二乗和を用いる場合は、加算器106の代わりに G_{in} と G_{local} に対する二乗和演算器を用いてもよい。すなわち、加算器106の代わりに G_{in} と G_{local} を入力とする任意の演算器を用いた場合も本発明を適用することが

出来る。

[0120] 以上の動作により、演算対象の学習ノード $2-0 \sim 2-1$, $2-3$ に勾配の和 ΣG が送信され、各学習ノード $2-0 \sim 2-1$, $2-3$ は、勾配の和 ΣG を用いてニューラルネットワークの構成パラメータを更新し、分散学習の 1 サイクルが終了する。

[0121] 以上説明したように、第 1 の実施の形態によれば、コンピューティングインタコネクタ装置 1 が各学習ノード 2 間の通信パケットの送受信処理と $All-reduce$ 処理とを同時並行して実行する。そのため、ヘッドノードで通信処理や $All-reduce$ 処理を実行する場合に比べて学習を高速化でき、通信ネットワーク 3 で接続された各学習ノード 2 間での協調処理をより高速に行うことができる。

[0122] また、第 1 の実施の形態によれば、各学習ノード 2 が対になって接続されたコンピューティングインタコネクタ装置 1 を通じてリング型の通信ネットワーク 3 に接続される構成であるので、接続される学習ノード 2 の数が増えた場合でも、リング型の通信ネットワーク 3 の通信帯域は、学習ノード 2 の数によらず一定でよいという利点がある。

[0123] [第 2 の実施の形態]

次に、本発明の第 2 の実施の形態について説明する。なお、以下の説明では、上述した第 1 の実施の形態と同じ構成については同一の符号を付し、その説明を省略する。

[0124] 第 1 の実施の形態では、コンピューティングインタコネクタ装置 1 が $All-reduce$ 処理のみを行う場合について説明した。これに対し、第 2 の実施の形態に係る分散深層学習システムでは、コンピューティングインタコネクタ装置 1' においてニューラルネットワークの構成パラメータの更新演算についてもさらに行う点で第 1 の実施の形態とは異なる。

[0125] [分散深層学習システムの動作]

図 17A～図 17G は、第 2 の実施の形態に係る分散深層学習システムの動作を説明する図である。

まず、図17Aに示すように、コンピューティングインタコネクタ装置1C1__0'に接続された学習ノード2-0から勾配の計算結果G0をコンピューティングインタコネクタ装置1C1__0'に送信する。そして、コンピューティングインタコネクタ装置1C1__0'は、勾配の計算結果G0をコンピューティングインタコネクタ装置1C1__1'に転送する。

[0126] 図17Bに示すように、コンピューティングインタコネクタ装置1C1__1'は、コンピューティングインタコネクタ装置1C1__0'から送信された勾配の計算結果G0と、コンピューティングインタコネクタ装置1C1__1'の直下の学習ノード2-1から送信された勾配の計算結果G1との和 $G0 + G1$ を計算する。コンピューティングインタコネクタ装置1C1__1'はこの計算結果 $G0 + G1$ を通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__2'に送信する。

[0127] 図17Cに示すように、本演算では、学習ノード2-2は演算対象外なので、コンピューティングインタコネクタ装置1C1__2'は、コンピューティングインタコネクタ装置1C1__1'から送信された勾配の計算結果 $G0 + G1$ に対して和演算を行わない。コンピューティングインタコネクタ装置1C1__2'は、計算結果 $G0 + G1$ をそのまま通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__3'へ送信する。

[0128] 図17Dに示すように、コンピューティングインタコネクタ装置1C1__3'では、コンピューティングインタコネクタ装置1C1__1'と同様の演算を行う。より詳細には、コンピューティングインタコネクタ装置1C1__3'は、コンピューティングインタコネクタ装置1C1__2'から送信された勾配の和の計算結果 $G0 + G1$ と、コンピューティングインタコネクタ装置1C1__3'の直下の学習ノード2-3から送信された勾配の計算結果G3との和 $\Sigma G = G0 + G1 + G3$ を計算する。

[0129] コンピューティングインタコネクタ装置1C1__3'は、この計算結果 ΣG を通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__0'に送信する。

[0130] また、図 1 7 D に示すように、勾配の和の計算結果 ΣG を受信したコンピューティングインタコネクト装置 1 C I __ 0' は、勾配の和 ΣG を用いてニューラルネットワークの構成パラメータの更新後の値 w_new を計算する。コンピューティングインタコネクト装置 1 C I __ 0' は、その計算結果をコンピューティングインタコネクト装置 1 C I __ 0' の直下の学習ノード 2-0 とコンピューティングインタコネクト装置 1 C I __ 1' とに通信ネットワーク 3 を介して送信する。

[0131] 図 1 7 E に示すように、構成パラメータの更新後の値 w_new を受信したコンピューティングインタコネクト装置 1 C I __ 1' は、構成パラメータの更新後の値 w_new をコンピューティングインタコネクト装置 1 C I __ 1' の直下の学習ノード 2-1 とコンピューティングインタコネクト装置 1 C I __ 2' とに送信する。

[0132] 図 1 7 F に示すように、構成パラメータの更新後の値 w_new を受信したコンピューティングインタコネクト装置 1 C I __ 2' は、構成パラメータの更新後の値 w_new を直下の学習ノード 2-2 へは送信せず、コンピューティングインタコネクト装置 1 C I __ 3' のみに通信ネットワーク 3 を介して送信する。

[0133] 図 1 7 F に示すように、コンピューティングインタコネクト装置 1 C I __ 3' は、コンピューティングインタコネクト装置 1 C I __ 2' から送信された構成パラメータの更新後の値 w_new をコンピューティングインタコネクト装置 1 C I __ 3' の直下の学習ノード 2-3 とコンピューティングインタコネクト装置 1 C I __ 0' とに送信する。

[0134] 最後に、図 1 7 G に示すように構成パラメータの更新後の値 w_new を受信したコンピューティングインタコネクト装置 1 C I __ 0' は、構成パラメータの更新後の値 w_new を廃棄する。

以上の動作により、演算対象の学習ノード 2-0 ~ 2 に構成パラメータの更新後の値 w_new が送信される。

[0135] [コンピューティングインタコネクト装置の構成]

次に、本実施の形態に係るコンピューティングインタコネクタ装置 1' の構成について、図 18 を参照して説明する。なお、本実施の形態に係る学習ノード 2 の構成は、第 1 の実施の形態と同様である。

[0136] コンピューティングインタコネクタ装置 1' は、ニューラルネットワーク (NN) 構成パラメータ更新演算部 110 と構成パラメータメモリ 111 とをさらに備える点以外は、第 1 の実施の形態に係るコンピューティングインタコネクタ装置 1 の構成 (図 9) と同様である。

[0137] NN 構成パラメータ更新演算部 110 は、ニューラルネットワークの構成パラメータの更新演算を行う。

更新パラメータメモリ 111 は、受信部 103 によってコンピューティングインタコネクタ装置 1' の直下に接続されている学習ノード 2 から受信された構成パラメータを記憶する。

[0138] [コンピューティングインタコネクタ装置の動作の具体例]

図 19 は、図 17A におけるコンピューティングインタコネクタ装置 1C1__0' の動作を示している。

学習ノード 2-0 から送信される通信パケット RPO のデータペイロードには、学習ノード 2-0 で計算された勾配値「G0」と、演算 ID「002」と、勾配値のシーケンシャル番号「003」と、受信完了フラグ「未了」とが格納されている。

[0139] コンピューティングインタコネクタ装置 1C1__0' の受信部 103 は、受信した通信パケット RPO のデータペイロードから勾配値 G0 と演算 ID とシーケンシャル番号と受信完了フラグとを取り出して振分部 104 に渡す。

[0140] 振分部 104 では、受信部 103 から受け取った演算 ID と演算情報テーブルをつきあわせて、この演算 ID = 2 に対しては、コンピューティングインタコネクタ装置 1C1__0' は「親」として動作すべきことを識別する。これにより、振分部 104 は、勾配値 G0 と演算 ID とシーケンシャル番号と受信完了フラグとを送信部 107 に渡す。

[0141] 送信部107は、振分部104から受け取った勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC1のデータペイロードに格納して、通信パケットTPC1を通信ネットワーク3を介して、隣接する下流側のコンピューティングインタコネクタ装置（図17Aの例ではコンピューティングインタコネクタ装置1C1__1'）へ送信する。

[0142] 図20は、図17Bにおけるコンピューティングインタコネクタ装置1C1__1'の動作を示している。

コンピューティングインタコネクタ装置1C1__1'の受信部100は、コンピューティングインタコネクタ装置1C1__0'から受信した通信パケットTPC1のデータペイロードから勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0143] 振分部101は、受信部100から受け取った受信完了フラグが「未了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクタ装置1C1__1'は「子」として動作すべきことを識別する。これにより、振分部101は、勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとをバッファメモリ102に格納する。

[0144] 一方、コンピューティングインタコネクタ装置1C1__1'の受信部103は、コンピューティングインタコネクタ装置1C1__1'の直下の学習ノード2-1から受信した通信パケットRP1のデータペイロードから勾配値G1と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部104に渡す。

[0145] 振分部104では、受信部103から受け取った演算IDと演算情報テーブルをつきあわせて、この演算ID=2に対しては、コンピューティングインタコネクタ装置1C1__1'は「子」として動作すべきことを識別する。これにより、振分部104は勾配値G0と演算IDとシーケンシャル番号と受信完了フラグとをバッファメモリ105に格納する。

[0146] コンピューティングインタコネクタ装置1C1__1'の制御部109は、

バッファメモリ 102 とバッファメモリ 105 に同一のシーケンシャル番号の勾配値「G0」と「G1」が揃った時点で、バッファメモリ 102 から勾配値 G0 とシーケンシャル番号と受信完了フラグとを読み出す。これと共に、制御部 109 は、バッファメモリ 105 から勾配値 G1 とシーケンシャル番号と受信完了フラグとを読み出し、勾配値「G0」と「G1」とを加算器に渡す。

[0147] 加算器 106 は、勾配値「G0」と「G1」とを加算する。また、制御部 109 は、バッファメモリ 102 から読み出した演算 ID とシーケンシャル番号と受信完了フラグとを送信部 107 に渡す。

[0148] コンピューティングインタコネクト装置 1C1__1' の送信部 107 は、加算器 106 によって計算された勾配の和「G0+G1」、および制御部 109 から受け取った演算 ID とシーケンシャル番号と受信完了フラグを通信パケット TPC2 のデータペイロードに格納する。そして、送信部 107 は、通信パケット TPC2 を、通信ネットワーク 3 を介してコンピューティングインタコネクト装置 1C1__1' に隣接する下流側のコンピューティングインタコネクト装置（図 17B の例ではコンピューティングインタコネクト装置 1C1__2'）へ送信する。

[0149] 図 21 は、図 17C におけるコンピューティングインタコネクト装置 1C1__2' の動作を示している。

コンピューティングインタコネクト装置 1C1__2' の受信部 100 は、コンピューティングインタコネクト装置 1C1__1' から受信した通信パケット TPC2 のデータペイロードから勾配値 G0+G1 と演算 ID とシーケンシャル番号と受信完了フラグとを取り出して振分部 101 に渡す。

[0150] 振分部 101 は、受け取った演算 ID と演算情報テーブルをつきあわせてこの演算 ID = 2 に対してはコンピューティングインタコネクト装置 1C1__2' は「演算対象外（非親子）」として動作すべきことを識別する。これにより、振分部 101 は、勾配値 G0+G1 と演算 ID とシーケンシャル番号と受信完了フラグとを送信部 107 に送信する。

[0151] 送信部107は、振分部101から受け取った勾配値 $G_0 + G_1$ と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC3のデータペイロードに格納する。送信部107は、通信パケットTPC3を、通信ネットワーク3を介してコンピューティングインタコネクタ装置1C1__2'に隣接する下流側のコンピューティングインタコネクタ装置（図17Cの例ではコンピューティングインタコネクタ装置1C1__3'）へ送信する。

[0152] 図22は、図17Dにおけるコンピューティングインタコネクタ装置1C1__0'の動作を示している。

コンピューティングインタコネクタ装置1C1__0'の受信部100は、コンピューティングインタコネクタ装置1C1__0'に隣接する上流側のコンピューティングインタコネクタ装置（図17Dの例ではコンピューティングインタコネクタ装置1C1__3'）から受信した通信パケットTPC0のペイロードから勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0153] 振分部101は、受信部100から受け取った受信完了フラグが「未了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクタ装置1C1__0'は「親」として動作すべきことを識別する。

[0154] これにより、振分部101は、勾配の和 ΣG と演算IDとシーケンシャル番号と受信完了フラグとをNN構成パラメータ更新演算部110に渡す。このとき、振分部101は、受信部100から受け取った受信完了フラグを、「未了」から「完了」を示す値に変更した上でNN構成パラメータ更新演算部110に渡す。

[0155] 一方、学習開始時点において、演算対象の学習ノード2-0, 2-1, 2-3のニューラルネットワークは、同じ構成パラメータの初期値が設定されている。この構成パラメータの初期値をコンピューティングインタコネクタ装置1C1__0'の構成パラメータメモリ111に記憶しておく。

[0156] NN構成パラメータ更新演算部110は、振分部101から受け取った勾

配の和 ΣG と、構成パラメータメモリ111に記憶されている構成パラメータの値 w_old とを基に、ニューラルネットワークの構成パラメータの更新後の値 w_new を構成パラメータ毎に計算する。

[0157] NN構成パラメータ更新演算部110は、この計算結果と振分部101から受け取った演算IDとシーケンシャル番号と受信完了フラグとを送信部107, 108に出力する。NN構成パラメータ更新演算部110は、更新方法として例えば、勾配降下法を用いる場合は式(3)のような計算を行う。

[0158] また、NN構成パラメータ更新演算部110は、構成パラメータの更新後の値 w_new を送信部107, 108に出力すると同時に、構成パラメータメモリ111に格納されている構成パラメータの値を、更新後の値 w_new によって上書きする。

[0159] 送信部107は、NN構成パラメータ更新演算部110から受け取った構成パラメータの更新後の値 w_new と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC1のデータペイロードに格納する。送信部107は、通信パケットTPC1を、通信ネットワーク3を介してコンピューティングインタコネクト装置1C1_0'に隣接する下流側のコンピューティングインタコネクト装置(図17Dの例ではコンピューティングインタコネクト装置1C1_1')へ送信する。

[0160] コンピューティングインタコネクト装置1C1_0'の送信部108は、NN構成パラメータ更新演算部110から受け取った構成パラメータの更新後の値 w_new と演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTP0のデータペイロードに格納して、通信パケットTP0を学習ノード2-0へ送信する。

[0161] 図23は、図17Eにおけるコンピューティングインタコネクト装置1C1_1'の動作を示している。

コンピューティングインタコネクト装置1C1_1'の受信部100は、コンピューティングインタコネクト装置1C1_0'から受信した通信パケットTPC1のデータペイロードから構成パラメータの更新後の値 w_new

wと演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0162] 振分部101は、受信部100から受け取った受信完了フラグが「完了」を示していること、また、受け取った演算IDと演算情報テーブルをつきあわせてこの演算ID=2に対してはコンピューティングインタコネクト装置1C1__1'は「子」として動作すべきことを識別する。これにより、振分部101は、構成パラメータの更新後の値w__newと演算IDとシーケンシャル番号と受信完了フラグとを送信部107、および送信部108に渡す。

[0163] 送信部107は、振分部101から受け取った構成パラメータの更新後の値w__newと演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTPC2のデータペイロードに格納する。そして、送信部107は、通信パケットTPC2を、通信ネットワーク3を介してコンピューティングインタコネクト装置1C1__1'に隣接する下流側のコンピューティングインタコネクト装置（図17Eの例ではコンピューティングインタコネクト装置1C1__2'）へ送信する。

[0164] 送信部108は、振分部101から受け取った構成パラメータの更新後の値w__newと演算IDとシーケンシャル番号と受信完了フラグとを通信パケットTP1のデータペイロードに格納して、通信パケットTP1を学習ノード2-1へ送信する。

[0165] 図24は、図17Fにおけるコンピューティングインタコネクト装置1C1__2'の動作を示している。

コンピューティングインタコネクト装置1C1__2'の受信部100は、コンピューティングインタコネクト装置1C1__1'から受信した通信パケットTPC2のデータペイロードから構成パラメータの更新後の値w__newと演算IDとシーケンシャル番号と受信完了フラグとを取り出して振分部101に渡す。

[0166] 振分部101は、受け取った演算IDと演算情報テーブルをつきあわせて

この演算 $ID = 2$ に対してはコンピューティングインタコネクタ装置 $1C1_2'$ は「演算対象外（非親子）」として動作すべきことを識別する。これにより、振分部 101 は、構成パラメータの更新後の値 w_new と演算 ID とシーケンシャル番号と受信完了フラグとを送信部 107 に送信する。

[0167] 送信部 107 は、振分部 101 から受け取った構成パラメータの更新後の値 w_new と演算 ID とシーケンシャル番号と受信完了フラグとを通信パケット $TPC3$ のデータペイロードに格納する。その後、送信部 107 は、通信パケット $TPC3$ を通信ネットワーク 3 を介してコンピューティングインタコネクタ装置 $1C1_2'$ に隣接する下流側のコンピューティングインタコネクタ装置（図 $17F$ の例ではコンピューティングインタコネクタ装置 $1C1_3'$ ）へ送信する。

[0168] 図 25 は、図 $17G$ におけるコンピューティングインタコネクタ装置 $1C1_0'$ の動作を示している。

コンピューティングインタコネクタ装置 $1C1_0'$ の受信部 100 は、コンピューティングインタコネクタ装置 $1C1_0'$ に隣接する上流側のコンピューティングインタコネクタ装置（図 $17G$ の例ではコンピューティングインタコネクタ装置 $1C1_3'$ ）から受信した通信パケット $TPC0$ のペイロードから構成パラメータの更新後の値 w_new とシーケンシャル番号と受信完了フラグとを取り出して振分部 101 に渡す。

[0169] 振分部 101 は、受信部 100 から受け取った受信完了フラグが「完了」を示していること、また、受け取った演算 ID と演算情報テーブルをつきあわせてこの演算 $ID = 2$ に対してはコンピューティングインタコネクタ装置 $1C1_0'$ は「親」として動作すべきことを識別する。その後、振分部 101 は、受信部 100 から受け取った構成パラメータの更新後の値 w_new と演算 ID とシーケンシャル番号と受信完了フラグとを廃棄する。

[0170] 以上の動作により、演算対象の学習ノード $2-0$ 、 $2-1$ 、 $2-3$ に構成パラメータの更新後の値 w_new が送信される。演算対象の学習ノード $2-0$ 、 $2-1$ 、 $2-3$ は、シーケンシャル番号で特定される、ニューラルネ

ットワーク 26 の構成パラメータを、構成パラメータの更新後の値 w_{new} によって上書きすることにより、ニューラルネットワーク 26 を更新する。

[0171] なお、ここでは、各勾配の和を用いて重みパラメータの更新処理を行う場合を例に説明したが、各勾配の和の代わりに各勾配の重み付き和を用いる場合は、第 1 の実施の形態と同様に加算器 106 の代わりに G_{in} と G_{local} に対する重み付き和演算器を用いてもよい。また、各勾配の和の代わりに各勾配の二乗和を用いる場合は、加算器 106 の代わりに G_{in} と G_{local} に対する二乗和演算器を用いてもよい。すなわち、加算器 106 の代わりに G_{in} と G_{local} を入力とする任意の演算器を用いた場合も本発明を適用することが出来る。

[0172] 以上説明したように、第 2 の実施の形態によれば、ニューラルネットワークの構成パラメータの更新演算処理を行う NN パラメータ更新演算部 110 による専用演算回路を備えるので、学習処理をより高速化することができる。また、勾配の和演算も、構成パラメータの更新演算も、学習ノード 2 が有するニューラルネットワーク 26 の構成に依らず、構成パラメータ毎に独立して同じ演算を行えばよい。そのため、学習ノード 2-0 ~ 2-3 が備えるニューラルネットワーク 26 の構成を変えた場合でも、コンピューティングインタコネクタ装置 1' の演算器には同じ専用演算回路を用いることができるというメリットがある。

[0173] また、第 2 の実施の形態によれば、コンピューティングインタコネクタ装置 1' が各学習ノード 2 間の通信パケットの送受信処理と All-reduce 処理とを同時並行して高速にハードウェア処理できる。そのため、従来技術のようにヘッドノードで通信処理や All-reduce 処理をソフトウェア処理する場合に比べて、学習を高速化でき、通信ネットワーク 3 で接続された各学習ノード 2 間での協調処理をより高速に行うことができる。

[0174] また、第 2 の実施の形態によれば、各学習ノード 2 が対になって接続されたコンピューティングインタコネクタ装置 1 を通じてリング型の通信ネット

ワーク 3 に接続される構成であるので、接続される学習ノード 2 の数が増えた場合でも、リング型の通信ネットワーク 3 の通信帯域は、学習ノード 2 の数によらず一定でよいという利点がある。

[0175] 以上、本発明の分散深層学習システム、分散深層学習方法、およびコンピューティングインタコネクト装置における実施の形態について説明したが、本発明は説明した実施の形態に限定されるものではなく、請求項に記載した発明の範囲において当業者が想定し得る各種の変形を行うことが可能である。

符号の説明

[0176] 1, 1', 1C1__0~1C1__3…コンピューティングインタコネクト装置、2, 2-0~2-3…学習ノード、3…通信ネットワーク、20…入力部、21…損失関数計算部、22…勾配計算部、23…送信部、24…受信部、25…構成パラメータ更新部、26…ニューラルネットワーク、100, 103…受信部、101, 104…振分部、102, 105…バッファメモリ、106…加算器、107, 108…送信部、109…制御部、110…NNパラメータ更新演算部、111…構成パラメータメモリ。

請求の範囲

[請求項1]

1 方向に通信可能なリング型の通信ネットワークを介して互いに接続された複数のコンピューティングインタコネクト装置と、

前記複数のコンピューティングインタコネクト装置のそれぞれと一対一に接続された複数の学習ノードとを備え、

各コンピューティングインタコネクト装置は、

自装置に接続された学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信部と、

自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信部と、

この第2受信部が受信した前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信部によって取得された前記転送データを振り分ける第1振分部と、

前記第1受信部が受信した前記パケットに含まれる前記受信完了フラグと前記役割とに応じて、前記第1受信部によって取得された前記ノードデータを振り分ける第2振分部と、

前記第2振分部によって振り分けられた前記ノードデータ、または前記第1振分部により振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信する第1送信部と、

前記第1振分部によって振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信部とを備え、

前記第1振分部は、

前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記第1送信部および前記

第2送信部に振り分け、

前記受信完了フラグがパケットの受信の完了を示し、かつ、前記役割が親である場合には、前記転送データを廃棄し、

前記第2振分部は、

前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記ノードデータを前記第1送信部に振り分け、

各学習ノードは、

学習データの入力に対して演算結果を出力するニューラルネットワークと、

データをパケット化して、自ノードと接続されたコンピューティングインタコネクタ装置に送信する第3送信部と、

自ノードと接続された前記コンピューティングインタコネクタ装置から送信されたパケットを受信して、このパケットに格納された前記転送データを取得する第3受信部と、

この第3受信部が取得した前記転送データに基づいて前記ニューラルネットワークの構成パラメータデータを更新する構成パラメータ更新部と

を備えることを特徴とする分散深層学習システム。

[請求項2]

請求項1に記載の分散深層学習システムにおいて、

前記コンピューティングインタコネクタ装置は、

前記第1振分部によって振り分けられた前記転送データと、前記第2振分部によって振り分けられた前記ノードデータとを入力とする演算を行う演算器をさらに備え、

前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記転送データを前記演算器に振り分け、

前記第2振分部は、前記受信完了フラグがパケットの受信の未了を

示し、かつ、前記役割が子である場合には、前記ノードデータを前記演算器に振り分け、

前記演算器は、前記転送データおよび前記ノードデータを入力とする演算結果を前記第1送信部に出力する

ことを特徴とする分散深層学習システム。

[請求項3]

請求項1または請求項2に記載の分散深層学習システムにおいて、前記コンピューティングインタコネクタ装置は、

前記ノードデータを記憶する機能をもつ構成パラメータメモリと、

前記第1振分部によって振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として、更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されたデータを更新する構成パラメータ更新演算部と、

をさらに備え、

前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記構成パラメータ更新演算部に振り分け、

前記構成パラメータ更新演算部は、算出した前記更新後の構成パラメータデータを前記第1送信部および前記第2送信部に出力し、

前記第1送信部は、前記更新後の構成パラメータデータをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクタ装置に送信し、

前記第2送信部は、前記更新後の構成パラメータデータをパケット化して、自装置と接続された前記学習ノードに送信する

ことを特徴とする分散深層学習システム。

[請求項4]

1方向に通信可能なリング型の通信ネットワークを介して互いに接続された複数のコンピューティングインタコネクタ装置と、

前記複数のコンピューティングインタコネクタ装置のそれぞれと一対一に接続された複数の学習ノードとを備える分散深層学習システム

における分散深層学習方法であって、

各コンピューティングインタコネクト装置が、

自装置に接続された学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信ステップと、

自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信ステップと、

この第2受信ステップで受信された前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信ステップで取得された前記転送データを振り分ける第1振分ステップと、

前記第1受信ステップで受信された前記パケットに含まれる前記受信完了フラグと前記役割とに応じて、前記第1受信ステップで取得された前記ノードデータを振り分ける第2振分ステップと、

前記第2振分ステップで振り分けられた前記ノードデータ、または前記第1振分ステップで振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信する第1送信ステップと、

前記第1振分ステップで振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信ステップとを備え、

各学習ノードが、

ニューラルネットワークに学習データを入力して演算結果を出力するニューラルネットワーク演算ステップと、

データをパケット化して、自ノードと接続された前記コンピューティングインタコネクト装置に送信する第3送信ステップと、

自ノードと接続された前記コンピューティングインタコネクト装置

から送信されたパケットを受信して、このパケットに格納された前記転送データを取得する第3受信ステップと、

この第3受信ステップで取得された前記転送データに基づいて前記ニューラルネットワークの構成パラメータデータを更新する構成パラメータ更新ステップと

を備えることを特徴とする分散深層学習方法。

[請求項5]

請求項4に記載の分散深層学習方法において、

前記コンピューティングインタコネクト装置が、

前記ノードデータを構成パラメータメモリに記憶する構成パラメータ記憶ステップと、

前記第1振分ステップで振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として、更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されたデータを更新する構成パラメータ更新演算ステップと、

をさらに備えることを特徴とする分散深層学習方法。

[請求項6]

1方向に通信可能なリング型の通信ネットワークを介して互いに接続され、かつ、複数の学習ノードのそれぞれと一対一に接続された複数のコンピューティングインタコネクト装置であって、

自装置に接続された学習ノードから送信されたパケットを受信して、このパケットに格納されたノードデータを取得する第1受信部と、

自装置に隣接する前記通信ネットワークの上流側のコンピューティングインタコネクト装置から送信されたパケットを受信して、このパケットに格納された転送データを取得する第2受信部と、

この第2受信部が受信した前記パケットに含まれるパケットの受信の完了または未了を示す受信完了フラグと自装置に対して予め割り当てられた役割とに応じて、前記第2受信部によって取得された前記転送データを振り分ける第1振分部と、

前記第1受信部が受信した前記パケットに含まれる前記受信完了フ

ラグと前記役割とに応じて、前記第1受信部によって取得された前記ノードデータを振り分ける第2振分部と、

前記第2振分部によって振り分けられた前記ノードデータ、または前記第1振分部により振り分けられた前記転送データをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクタ装置に送信する第1送信部と、

前記第1振分部によって振り分けられた前記転送データをパケット化して自装置と接続された前記学習ノードに送信する第2送信部とを備え、

前記第1振分部は、

前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記第1送信部および前記第2送信部に振り分け、

前記受信完了フラグがパケットの受信の完了を示し、かつ、前記役割が親である場合には、前記転送データを廃棄し、

前記第2振分部は、

記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記ノードデータを前記第1送信部に振り分ける

ことを特徴とするコンピューティングインタコネクタ装置。

[請求項7]

請求項6に記載のコンピューティングインタコネクタ装置において、

前記第1振分部によって振り分けられた前記転送データと、前記第2振分部によって振り分けられた前記ノードデータとを入力とする演算を行う演算器をさらに備え、

前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記転送データを前記演算器に振り分け、

前記第2振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が子である場合には、前記ノードデータを前記演算器に振り分け、

前記演算器は、前記転送データおよび前記ノードデータを入力とする演算結果を前記第1送信部に出力する

ことを特徴とするコンピューティングインタコネクト装置。

[請求項8]

請求項6または請求項7に記載のコンピューティングインタコネクト装置において、

前記ノードデータを記憶する機能をもつ構成パラメータメモリと、

前記第1振分部によって振り分けられた前記転送データと前記構成パラメータメモリに記憶されたデータを入力として更新後の構成パラメータデータを計算して、前記構成パラメータメモリに記憶されているデータを更新する構成パラメータ更新演算部と

をさらに備え、

前記第1振分部は、前記受信完了フラグがパケットの受信の未了を示し、かつ、前記役割が親である場合には、前記転送データを前記構成パラメータ更新演算部に振り分け、

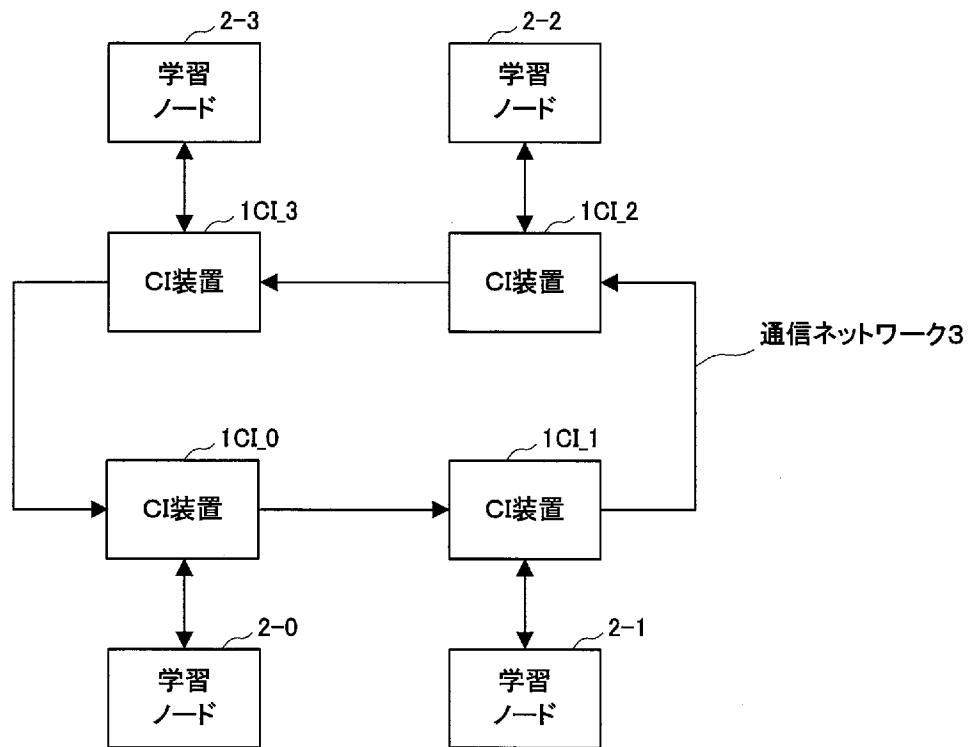
前記構成パラメータ更新演算部は、算出した前記更新後の構成パラメータデータを前記第1送信部および前記第2送信部に出力し、

前記第1送信部は、前記更新後の構成パラメータデータをパケット化して、自装置に隣接する前記通信ネットワークの下流側のコンピューティングインタコネクト装置に送信し、

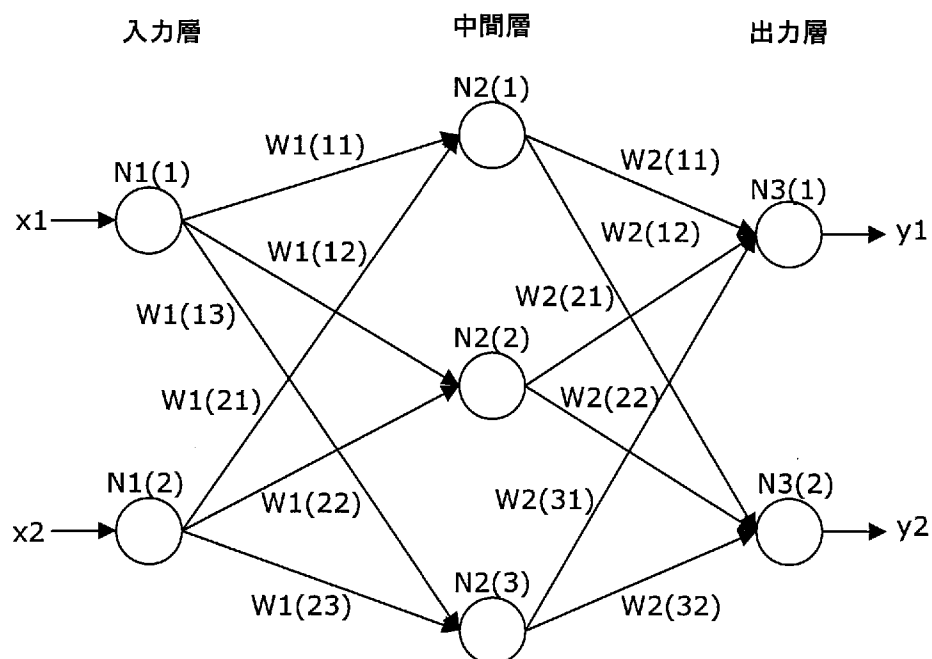
前記第2送信部は、前記更新後の構成パラメータデータをパケット化して、自装置と接続された前記学習ノードに送信する

ことを特徴とするコンピューティングインタコネクト装置。

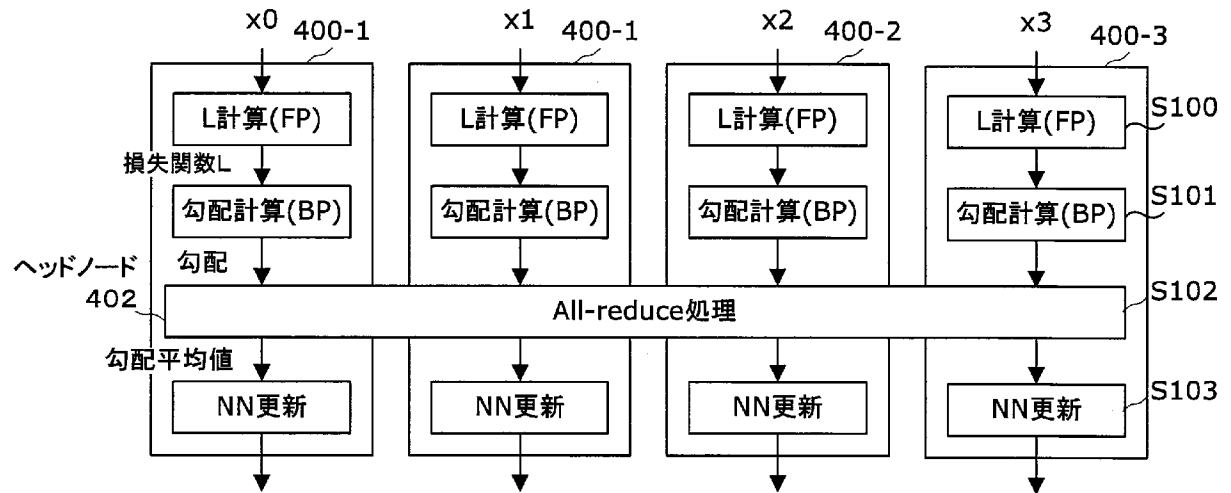
[図1]



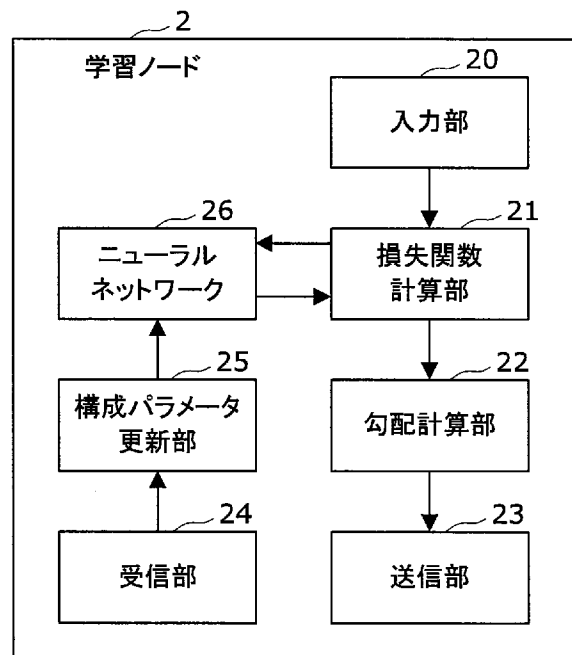
[図2]



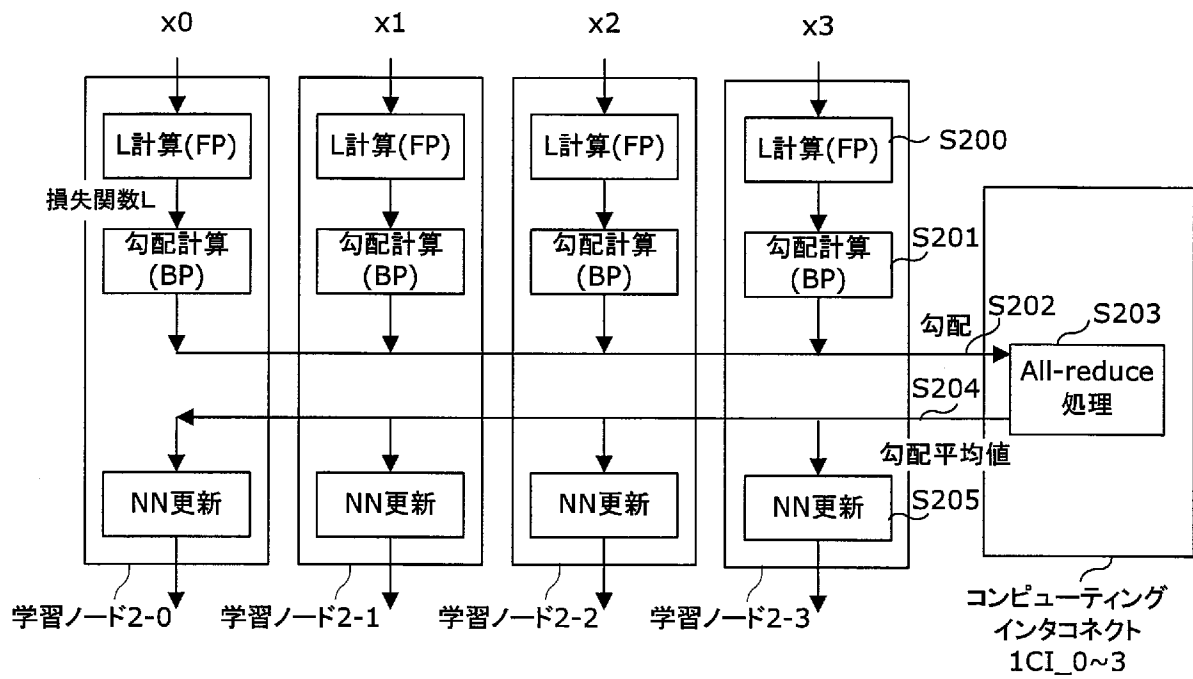
[図3]



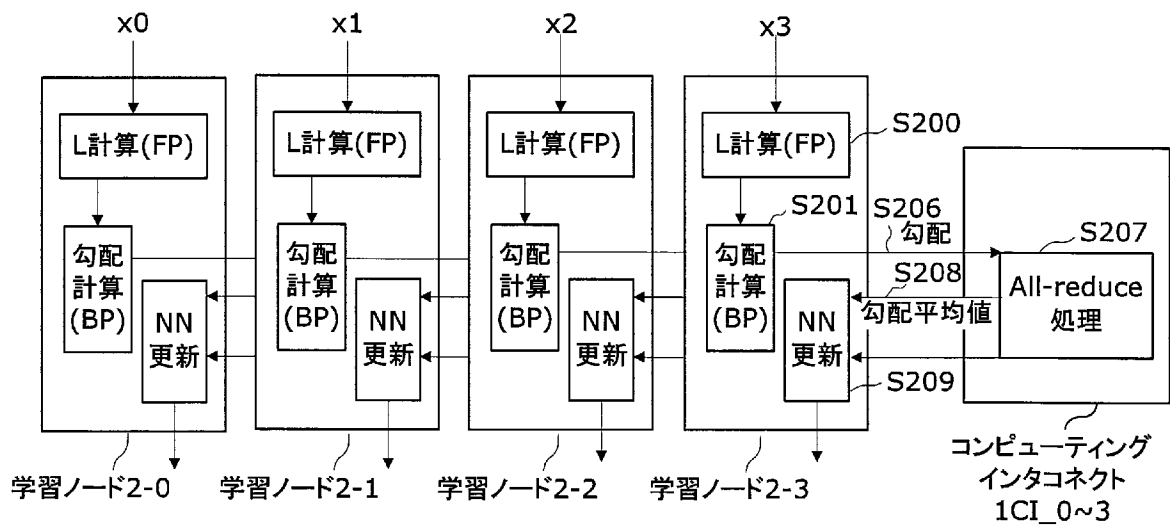
[図4]



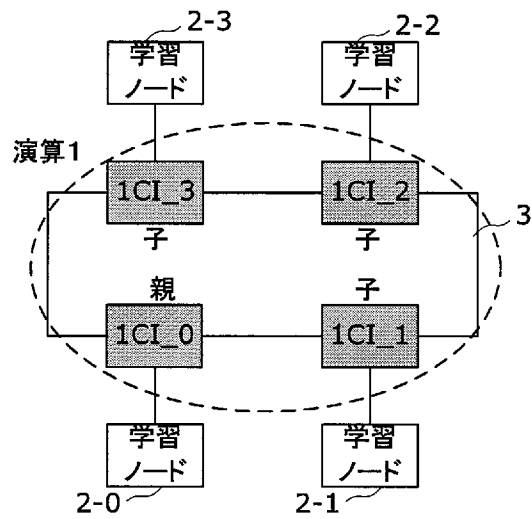
[図5]



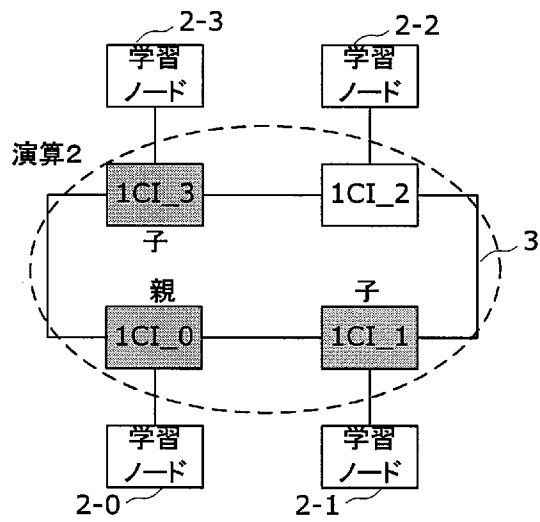
[図6]



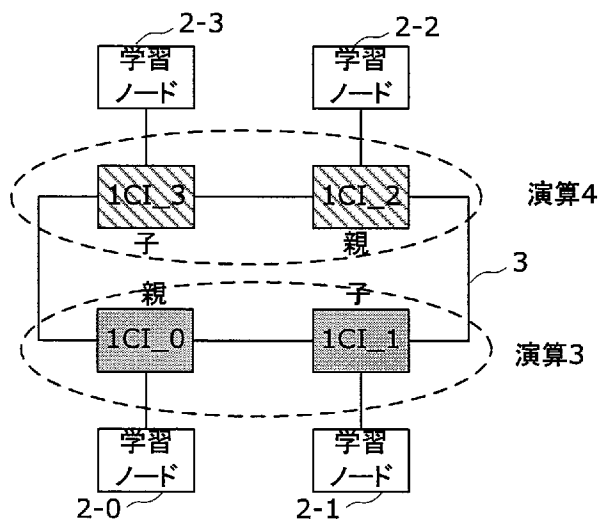
[図7A]



[図7B]



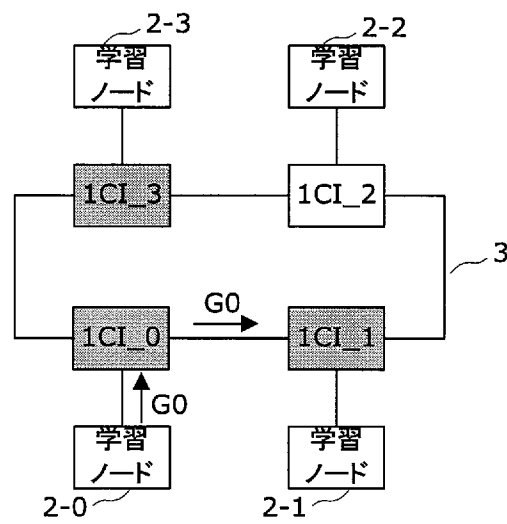
[図7C]



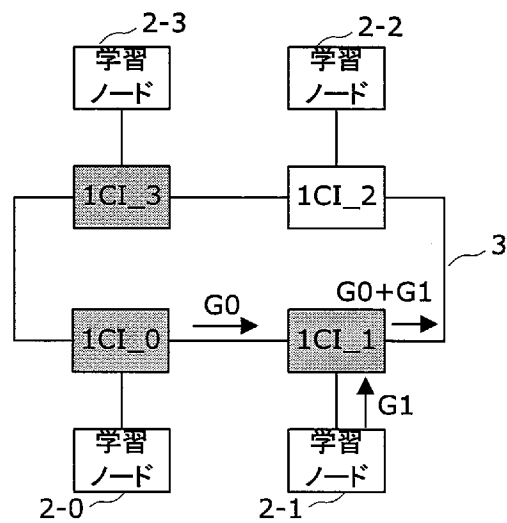
[図7D]

演算ID	1CI_0	1CI_1	1CI_2	1CI_3
1	親	子	子	子
2	親	子	対象外	子
3	親	子	対象外	対象外
4	対象外	対象外	親	子
...				

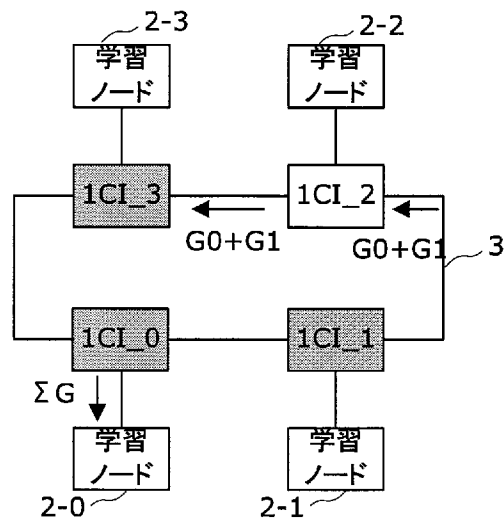
[図8A]



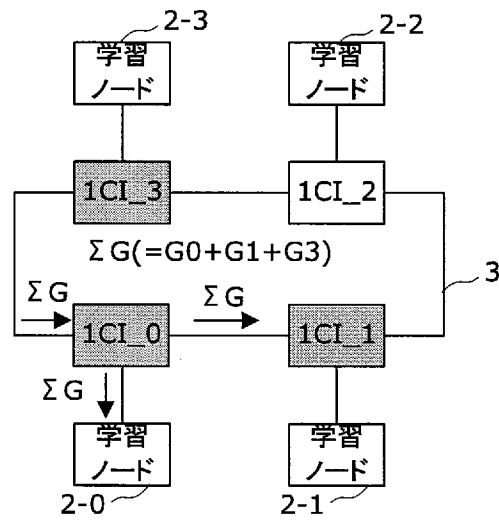
[図8B]



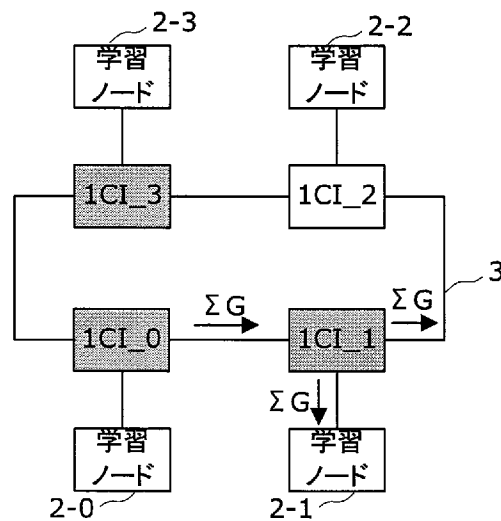
[図8C]



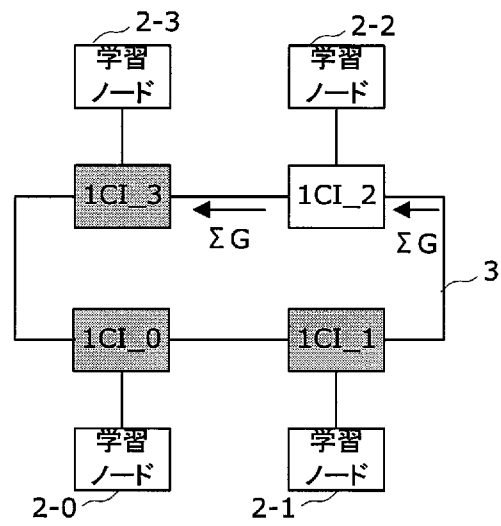
[図8D]



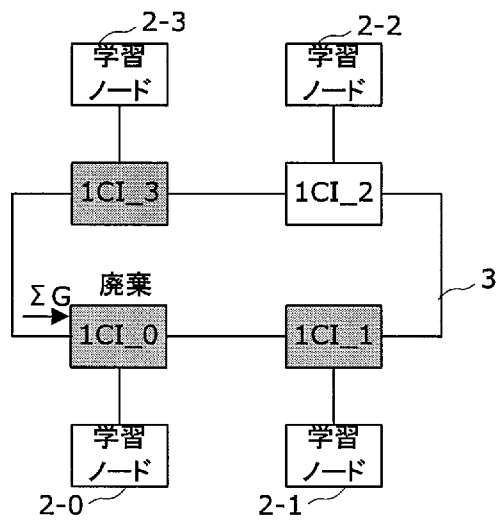
[図8E]



[図8F]



[図8G]



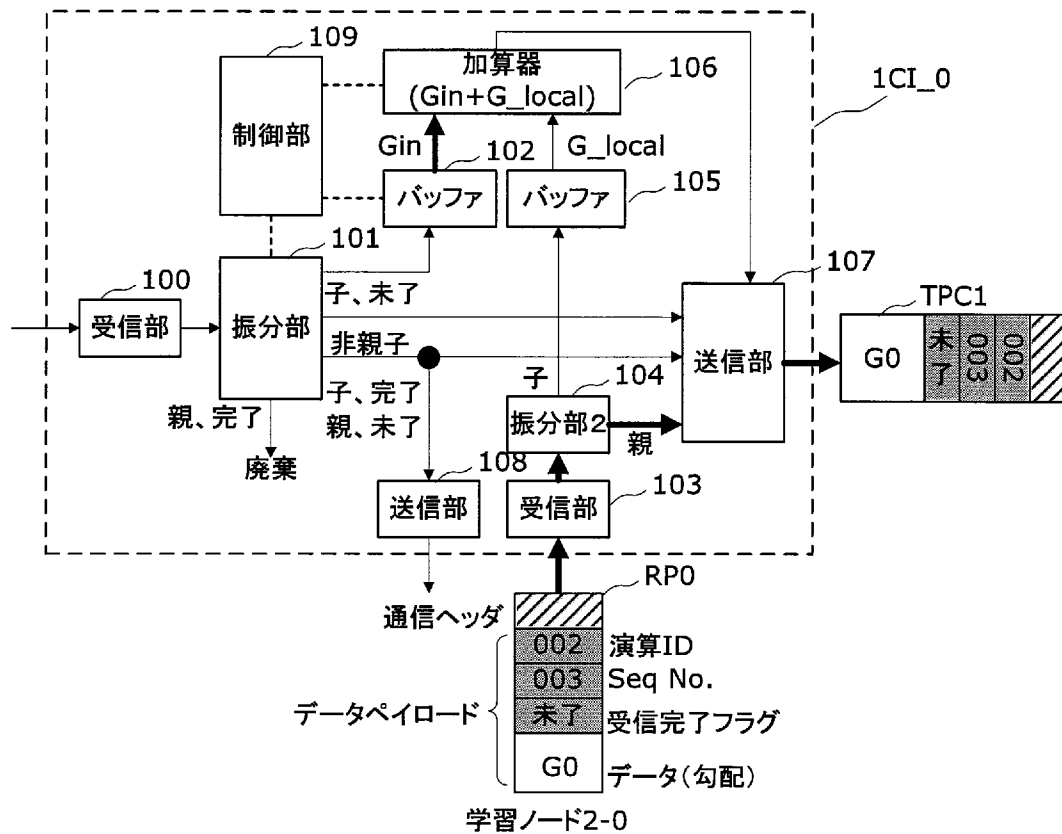
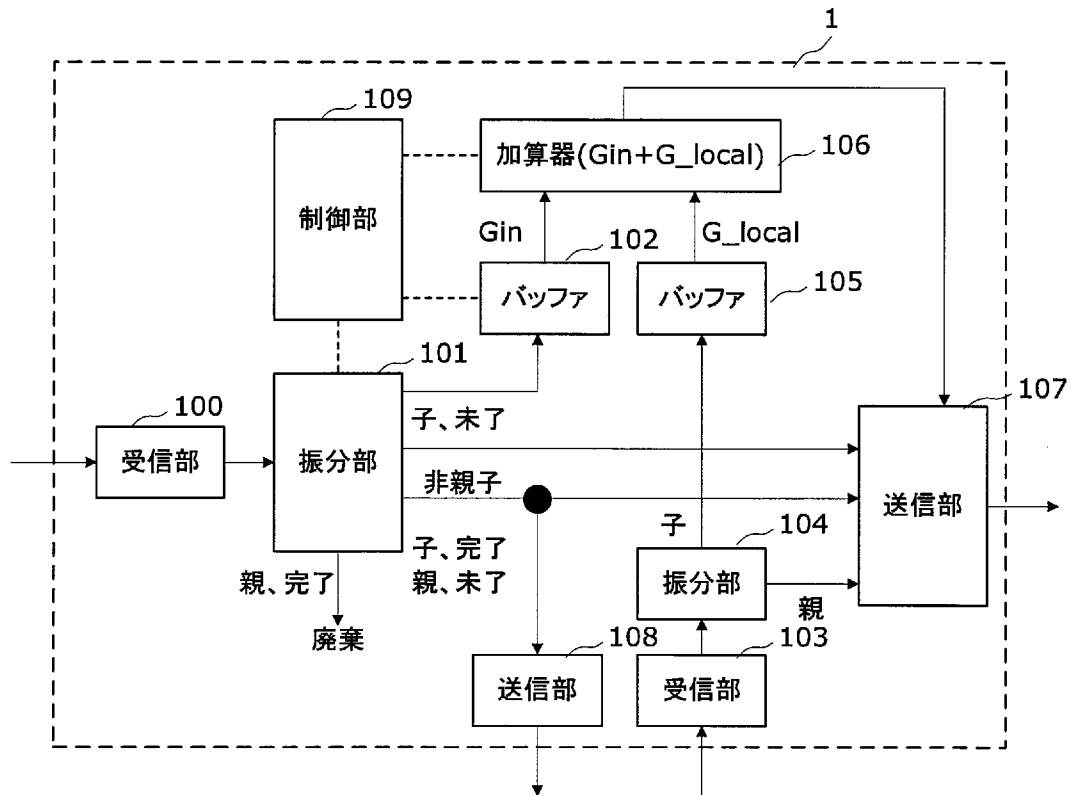
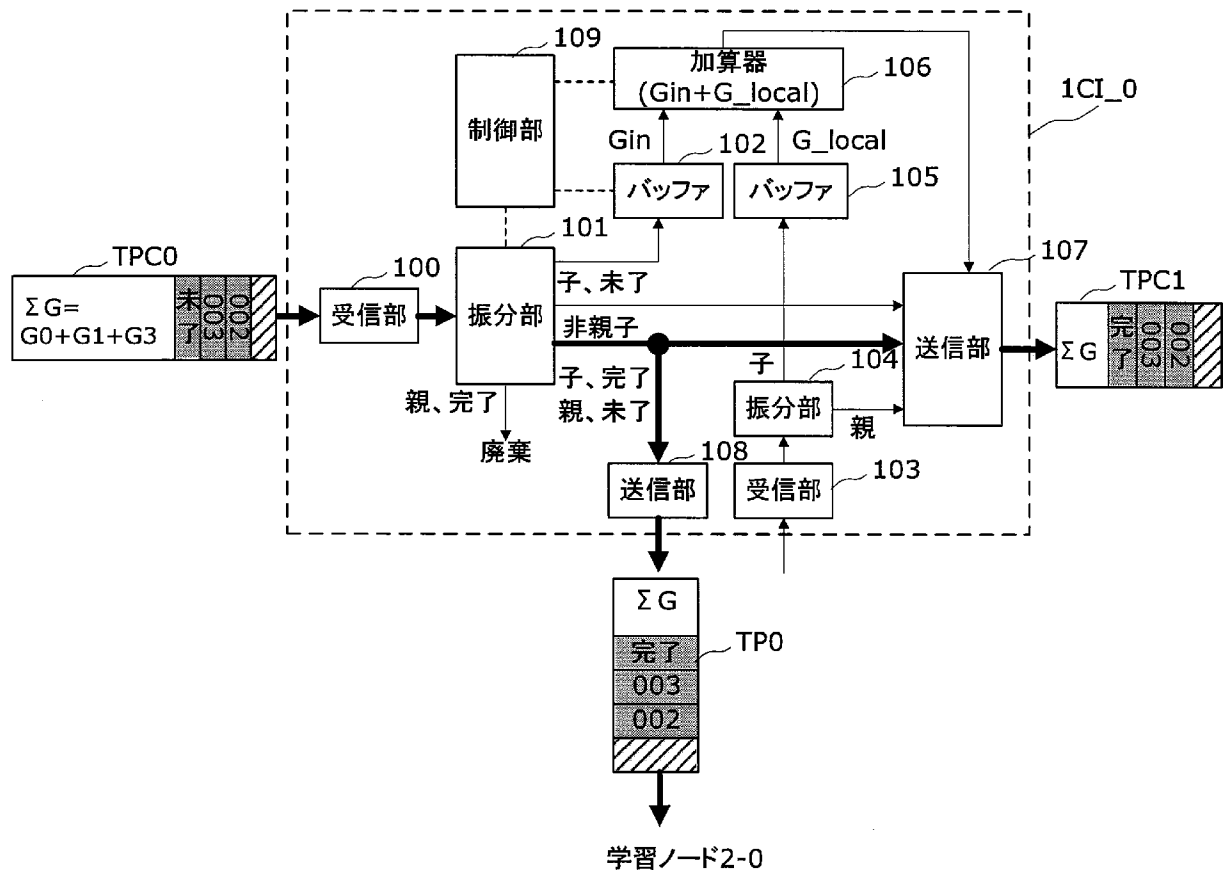


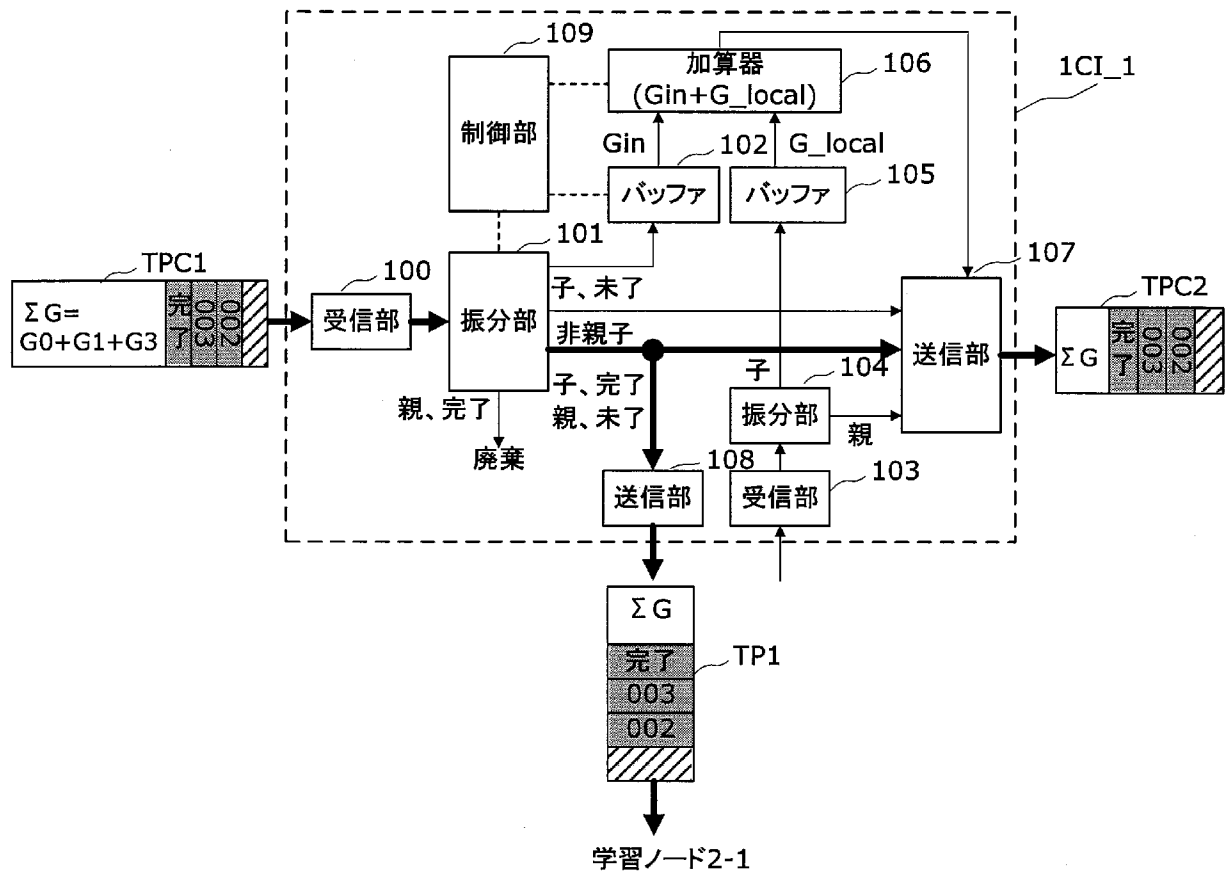
Figure 1 illustrates the first CI (1CI₁) and the data packet structure. The CI block contains a control unit (109), a buffer (102), an adder (106), and a transmission/reception unit (107). The data packet structure shows a communication header (RP1) with fields for operation ID (002), sequence number (003), reception completion flag (未了), and data (G1).

[illegible]

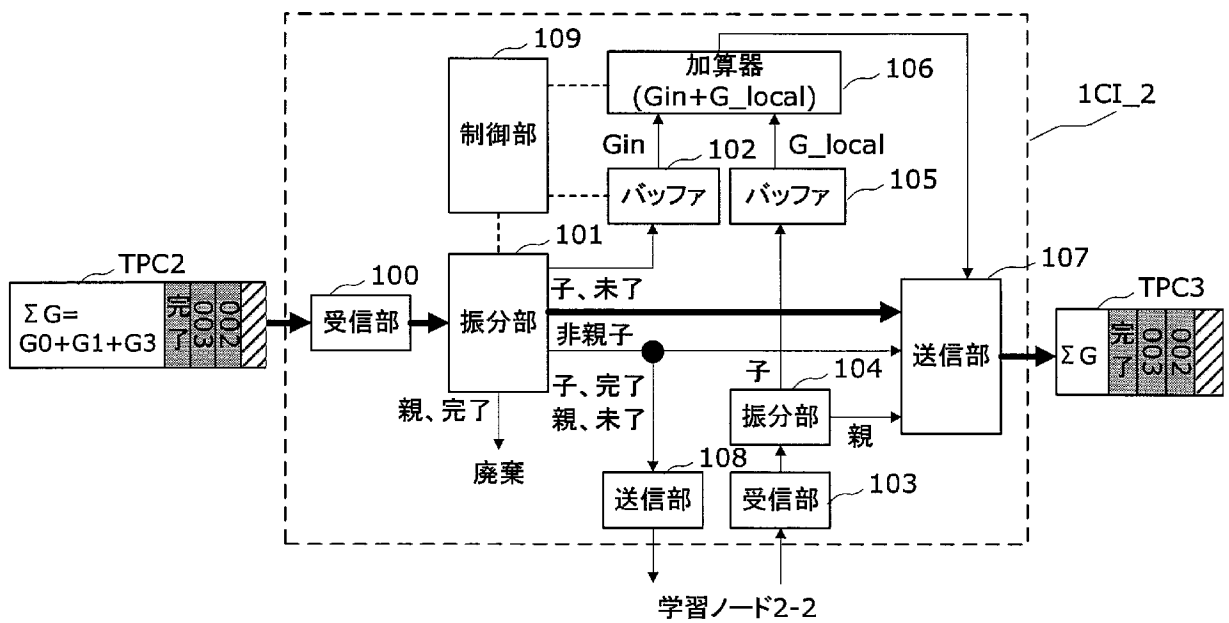
[図13]



[図14]



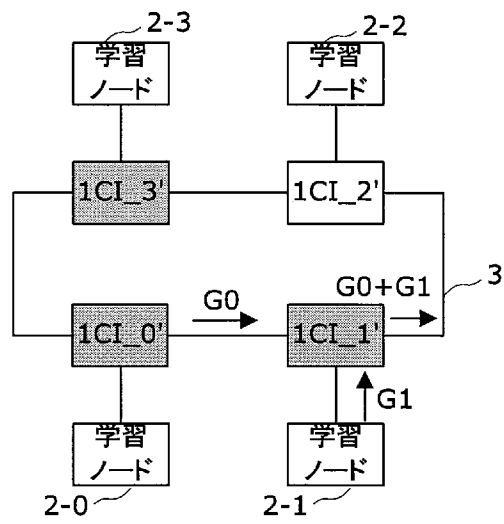
[図15]



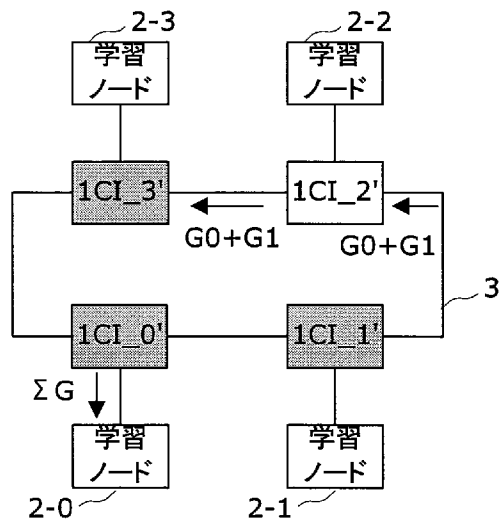
[illegible]

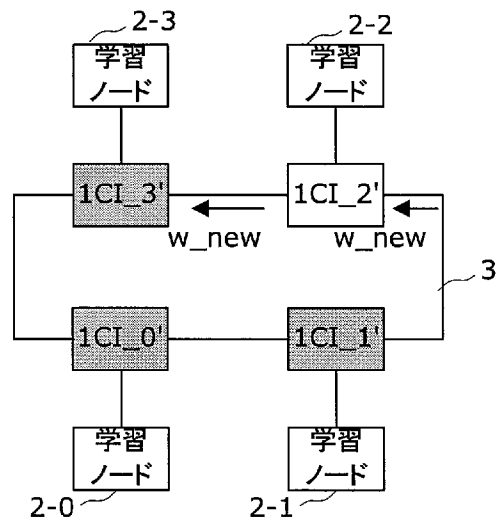
Diagram illustrating a network topology (FIG. 1). The network consists of four learning nodes (2-0, 2-1, 2-2, 2-3) and four intermediate nodes (1CI_0', 1CI_1', 1CI_2', 1CI_3'). The intermediate nodes 1CI_0' and 1CI_1' are shaded. The connections are as follows: Learning node 2-0 is connected to 1CI_0' via an arrow labeled 'G0'. Learning node 2-1 is connected to 1CI_1'. Learning node 2-2 is connected to 1CI_2'. Learning node 2-3 is connected to 1CI_3'. A line labeled '3' connects 1CI_1' to 1CI_2' and 1CI_3'. An arrow labeled 'G0' also points from 1CI_0' to 1CI_1'.

[図17B]

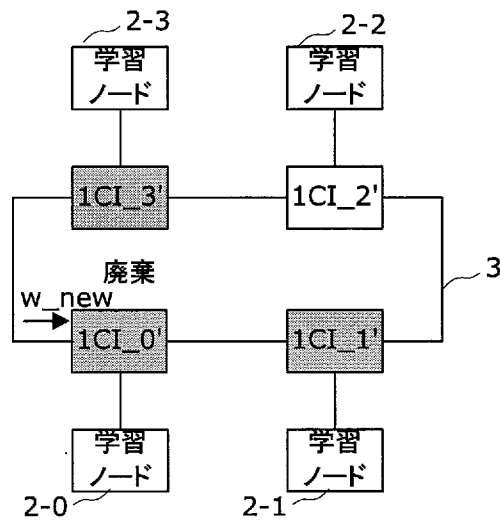


[図17C]

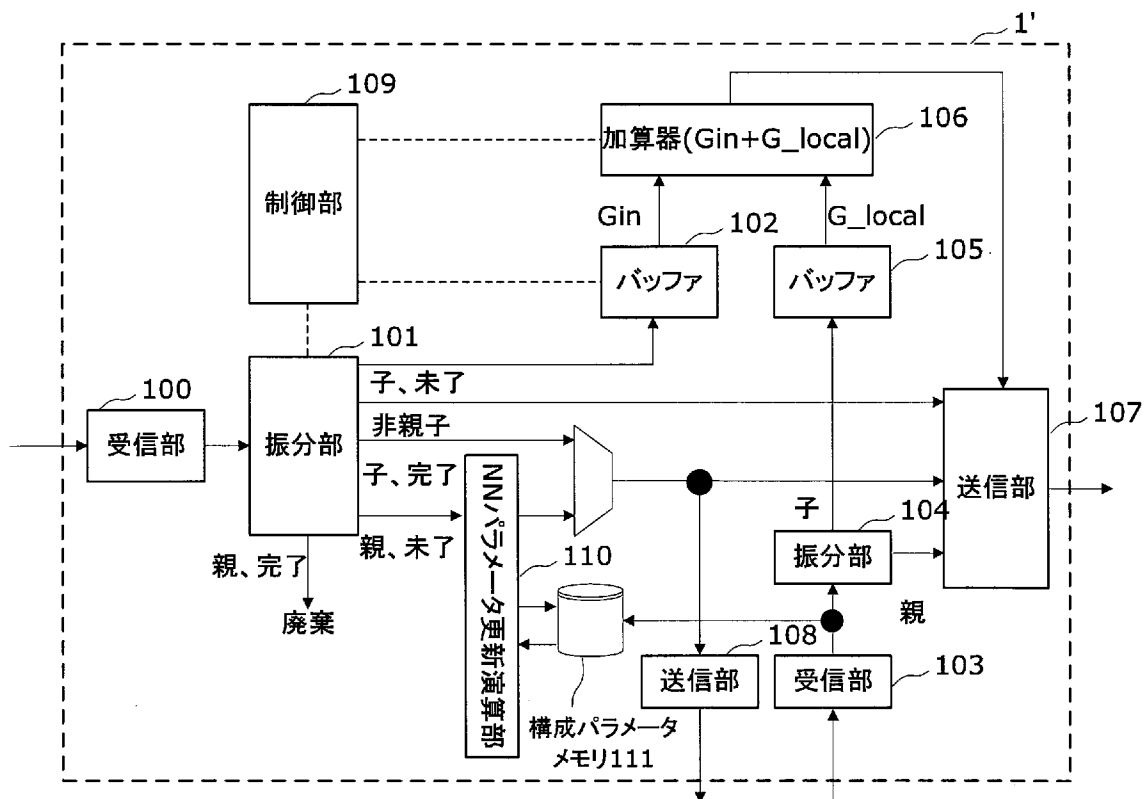




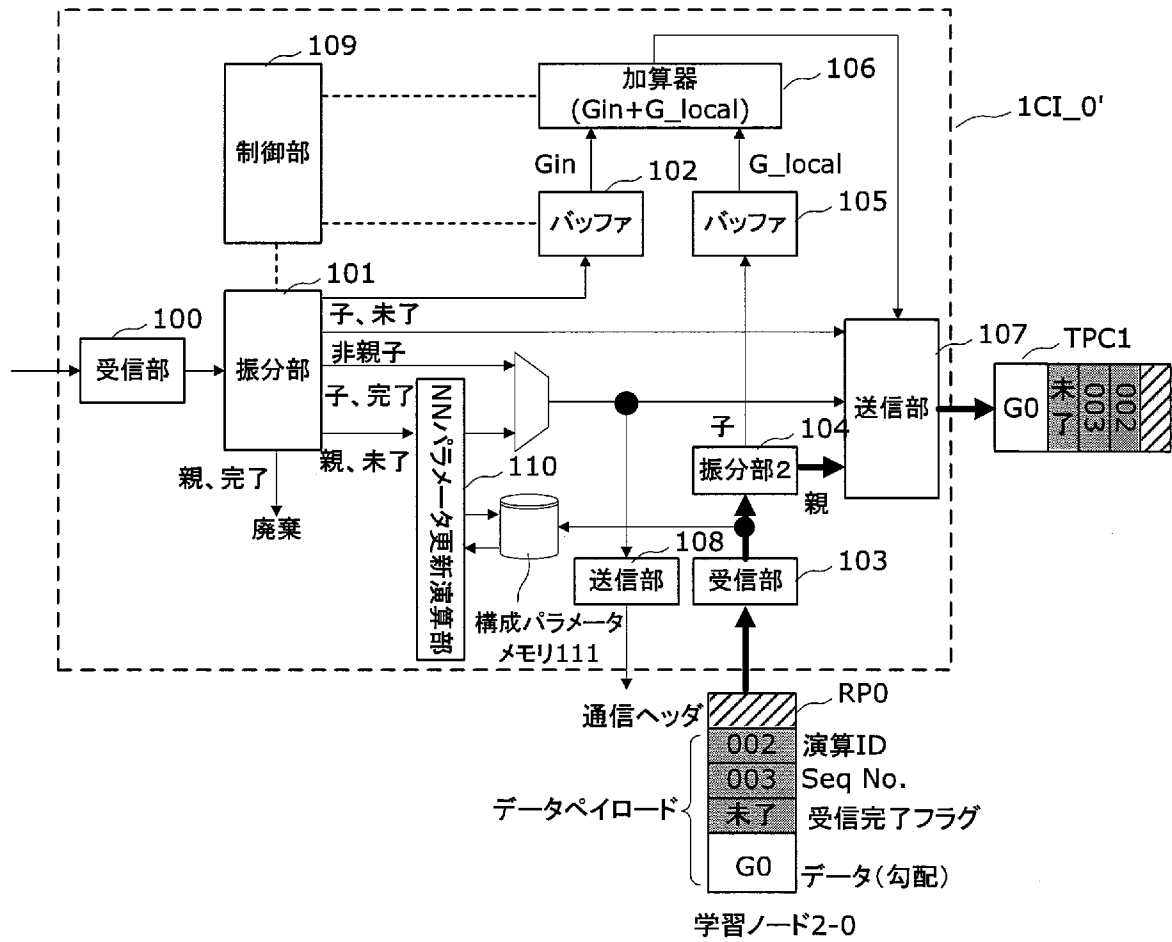
[図17G]



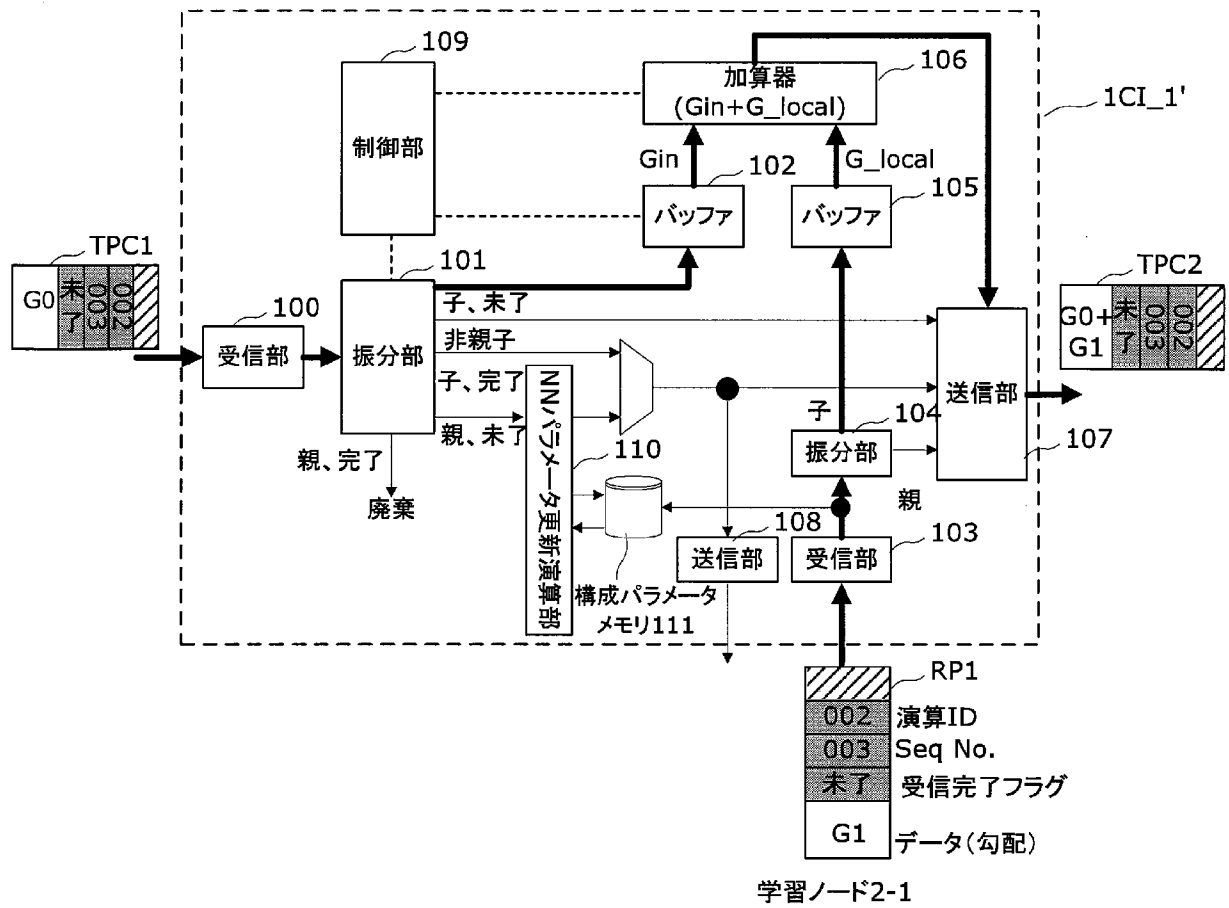
[図18]



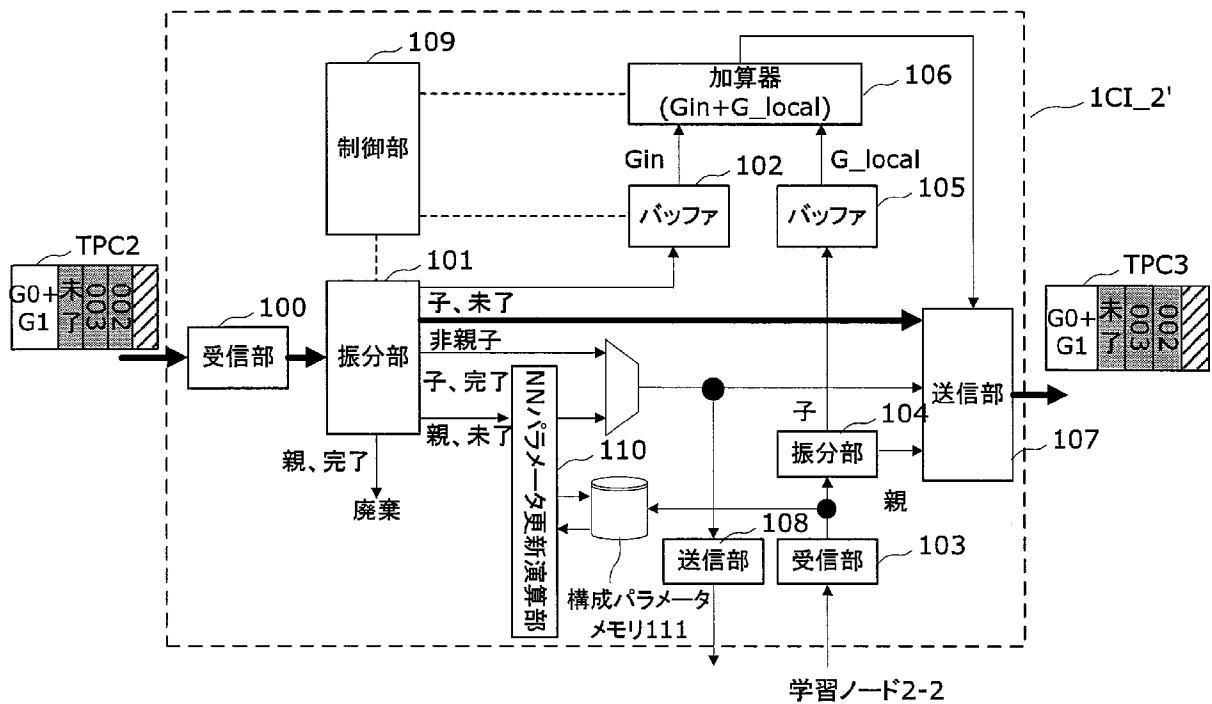
[図19]



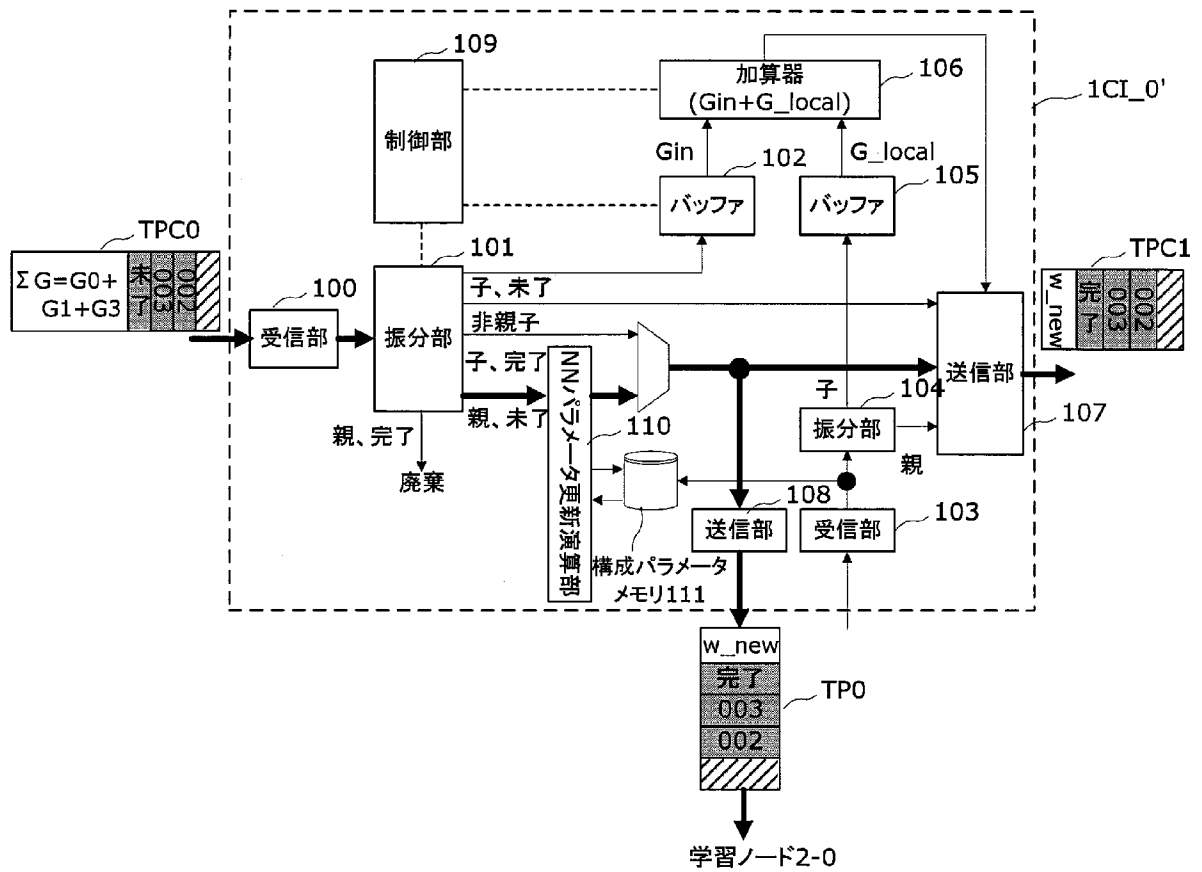
[図20]



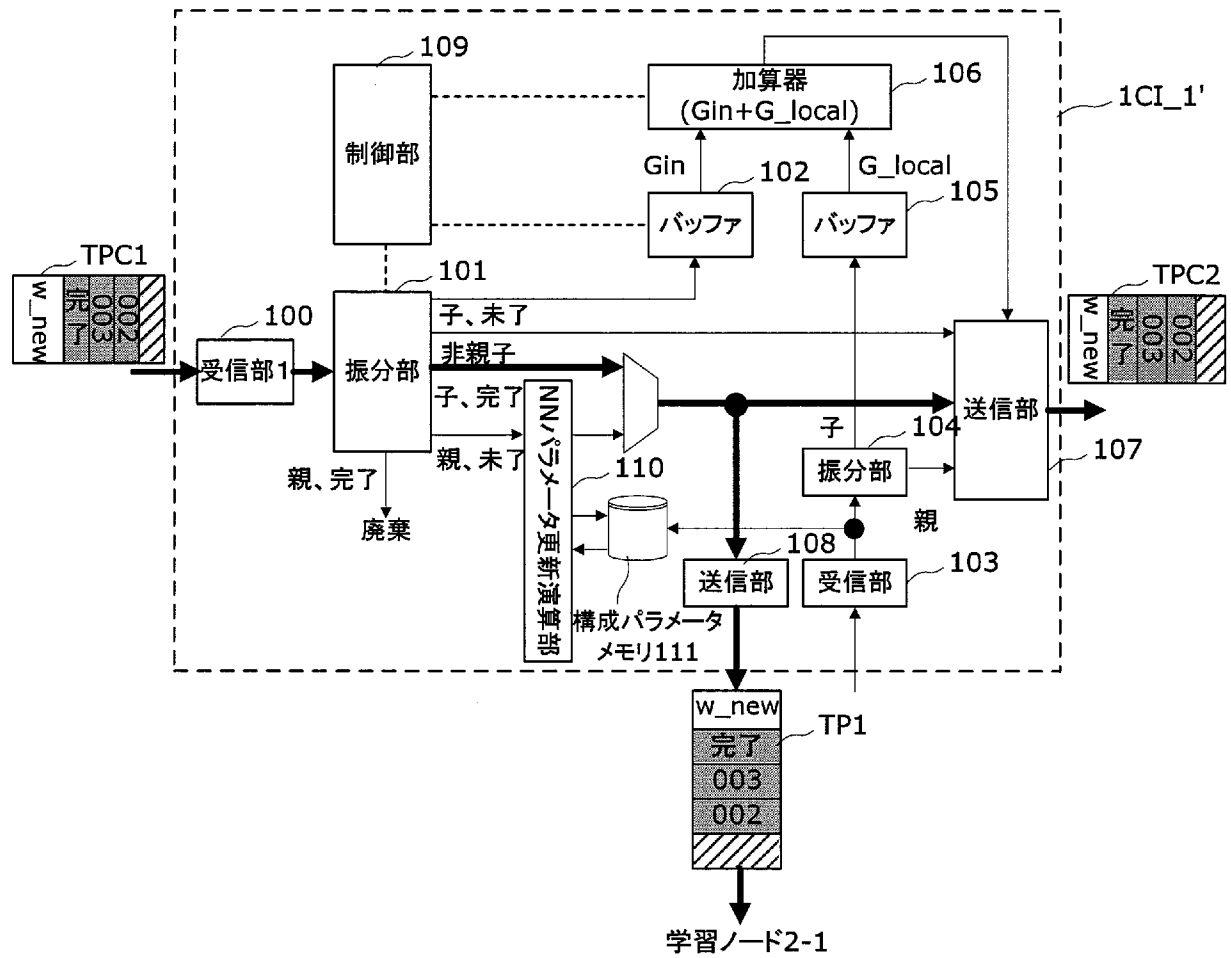
[図21]



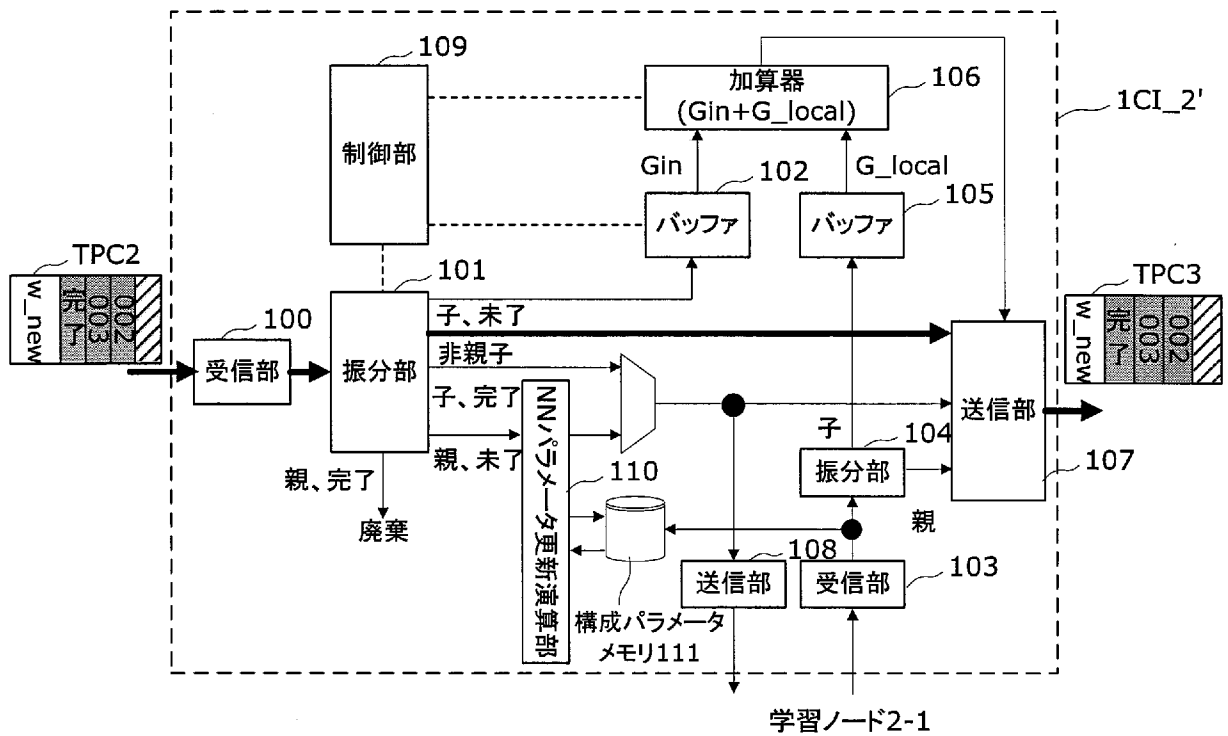
[図22]



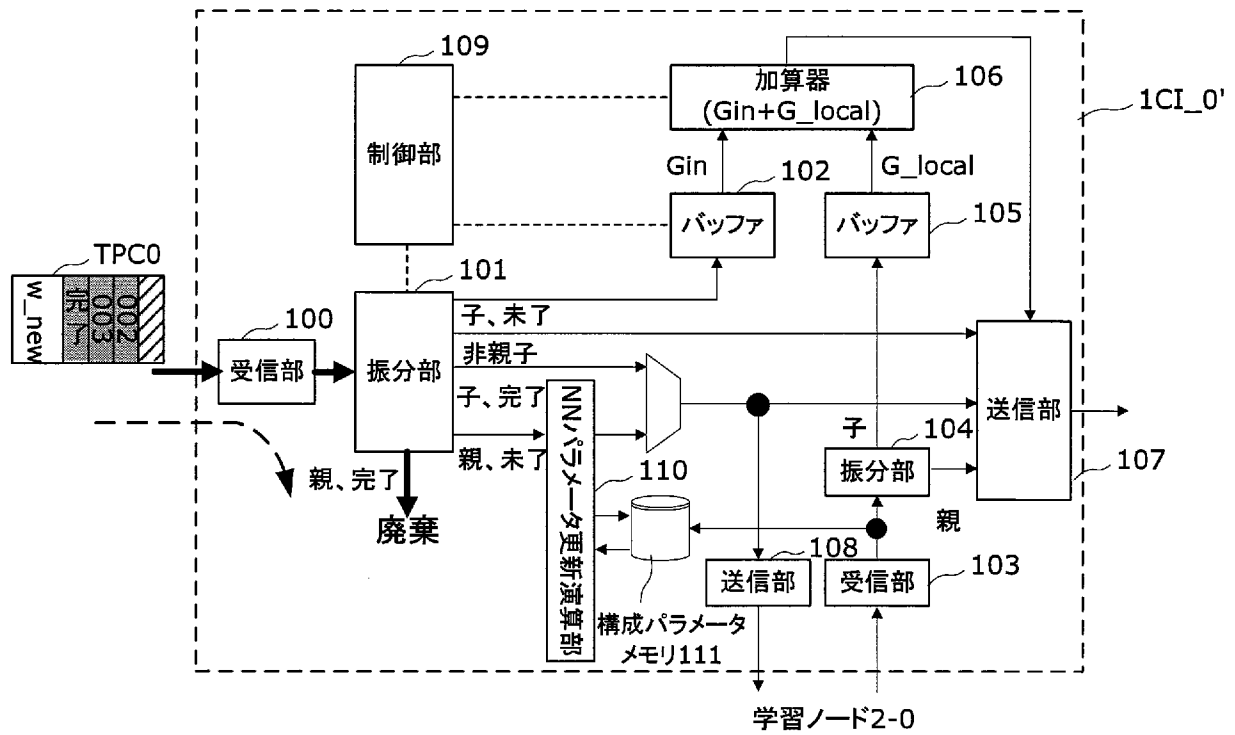
[図23]



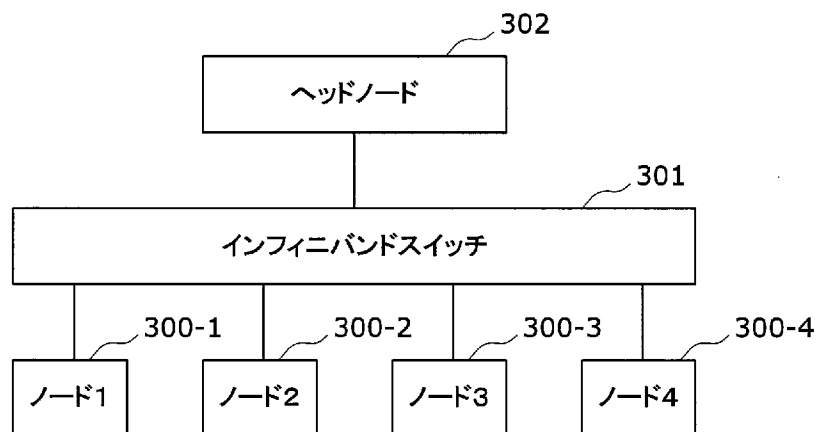
[図24]



[図25]



[図26]



INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2019/020906

A. CLASSIFICATION OF SUBJECT MATTER

Int.Cl. G06N3/08 (2006.01) i

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

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

Int.Cl. G06N3/08

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

Published examined utility model applications of Japan 1922-1996

Published unexamined utility model applications of Japan 1971-2019

Registered utility model specifications of Japan 1996-2019

Published registered utility model applications of Japan 1994-2019

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

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	JP 2018-36779 A (TOSHIBA CORP.) 08 March 2018, entire text (Family: none)	1-8
A	JP 2017-168950 A (HITACHI, LTD.) 21 September 2017, entire text (Family: none)	1-8

☐ Further documents are listed in the continuation of Box C.

☐ See patent family annex.

* Special categories of cited documents:

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

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

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

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

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

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

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

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

"&" document member of the same patent family

Date of the actual completion of the international search
09 July 2019 (09.07.2019)

Date of mailing of the international search report
23 July 2019 (23.07.2019)

Name and mailing address of the ISA/
Japan Patent Office
3-4-3, Kasumigaseki, Chiyoda-ku,
Tokyo 100-8915, Japan

Authorized officer

Telephone No.

A. 発明の属する分野の分類 (国際特許分類 (IPC))

Int.Cl. G06N3/08 (2006.01) i

B. 調査を行った分野

調査を行った最小限資料 (国際特許分類 (IPC))

Int.Cl. G06N3/08

最小限資料以外の資料で調査を行った分野に含まれるもの

日本国実用新案公報	1922-1996年
日本国公開実用新案公報	1971-2019年
日本国実用新案登録公報	1996-2019年
日本国登録実用新案公報	1994-2019年

国際調査で使用した電子データベース (データベースの名称、調査に使用した用語)

C. 関連すると認められる文献

引用文献の カテゴリー*	引用文献名 及び一部の箇所が関連するときは、その関連する箇所の表示	関連する 請求項の番号
A	JP 2018-36779 A (株式会社東芝) 2018.03.08, 全文 (ファミリーなし)	1-8
A	JP 2017-168950 A (株式会社日立製作所) 2017.09.21, 全文 (ファミリーなし)	1-8

C欄の続きにも文献が列挙されている。

パテントファミリーに関する別紙を参照。

* 引用文献のカテゴリー

「A」 特に関連のある文献ではなく、一般的技術水準を示すもの	「T」 国際出願日又は優先日後に公表された文献であって出願と矛盾するものではなく、発明の原理又は理論の理解のために引用するもの
「E」 国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの	「X」 特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの
「L」 優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献 (理由を付す)	「Y」 特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの
「O」 口頭による開示、使用、展示等に言及する文献	「&」 同一パテントファミリー文献
「P」 国際出願日前で、かつ優先権の主張の基礎となる出願	

国際調査を完了した日

09.07.2019

国際調査報告の発送日

23.07.2019

国際調査機関の名称及びあて先

日本国特許庁 (ISA/J P)

郵便番号 100-8915

東京都千代田区霞が関三丁目4番3号

特許庁審査官 (権限のある職員)

今城 朋彬

5 B

7888

電話番号 03-3581-1101 内線 3545