

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号
特許第7547526号
(P7547526)

(45)発行日 令和6年9月9日(2024.9.9)

(24)登録日 令和6年8月30日(2024.8.30)

(51)国際特許分類		F I	
G 0 6 F	12/0875(2016.01)	G 0 6 F	12/0875 1 0 6
G 0 6 F	9/38 (2018.01)	G 0 6 F	9/38 3 7 0 X
G 0 6 F	12/0806(2016.01)	G 0 6 F	12/0806 1 0 0
G 0 6 F	13/38 (2006.01)	G 0 6 F	13/38 3 4 0 C
請求項の数 6 外国語出願 (全10頁)			
(21)出願番号	特願2023-33382(P2023-33382)	(73)特許権者	516291125
(22)出願日	令和5年3月6日(2023.3.6)		レール ビジョン リミテッド
(62)分割の表示	特願2020-512508(P2020-512508))の分割		R A I L V I S I O N L T D
原出願日	平成30年8月30日(2018.8.30)		イスラエル国 4 3 6 6 5 1 7 ラアナナ
(65)公開番号	特開2023-78204(P2023-78204A)		ピー. オー. ボックス 2 1 5 5 レホヴ
(43)公開日	令和5年6月6日(2023.6.6)	(74)代理人	ハティダール 1 5
審査請求日	令和5年3月15日(2023.3.15)		100120891
(31)優先権主張番号	62/552,475	(74)代理人	弁理士 林 一好
(32)優先日	平成29年8月31日(2017.8.31)		100165157
(33)優先権主張国・地域又は機関	米国(US)	(74)代理人	弁理士 芝 哲央
			100205659
		(74)代理人	弁理士 齋藤 拓也
			100126000
		(74)代理人	弁理士 岩池 満
			100185269
		最終頁に続く	

(54)【発明の名称】 複数計算における高スループットのためのシステムおよび方法

(57)【特許請求の範囲】

【請求項1】

データのストリームを受信するように構成され、前記データのストリームを、データ部分に分割し、前記データ部分の各々を、複数の先入れ先出し（F I F O）レジスタの1つを通過させるように構成された、複数の前記F I F Oレジスタを備えたアレイと、

前記データ部分を受信するように構成された、第1のアドバンスト・イクステンシブル・インタフェース（A X I）ユニットと、

を備えたデータ・ストリーム・デバイダ・ユニット（D S D U）（3 0 4）と、

前記第1のA X Iユニットからデータ部分を受信するように構成された、第2のアドバンスト・イクステンシブル・インタフェース（A X I）ユニットと、

それぞれのF I F Oレジスタから各データ部分を受信するように構成され、前記受信したデータ部分処理するように構成された、複数のストリーミングマルチプロセッサ（S M）と、

を備えたグラフィックプロセッシングユニット（G P U）と、

を備えた、G P Uでの処理がまだ施されていない未処理データを処理する回路。

【請求項2】

前記D S D U内のF I F Oレジスタは、前記D S D U内の、前記F I F Oレジスタに割り当てられた第1のA X Iユニットと、前記G P U内の、前記F I F Oレジスタに割り当てられた第2のA X Iユニットを介して前記G P U内の、前記F I F Oレジスタに割り当てられたS Mに接続される、請求項1に記載の回路。

【請求項 3】

前記 D S D U 内の、前記 F I F O レジスタの各々は、前記 D S D U 内の、第 1 の共通 A X I ユニット、および前記 G P U 内の、共通 A X I ユニットの介して、前記 G P U 内の、前記 F I F O レジスタに割り当てられた S M に接続される、請求項 1 に記載の回路。

【請求項 4】

グラフィックプロセッサユニット (G P U) での処理がまだ施されていない未処理データのストリームを受信するステップと、

前記ストリームを複数のデータ部分に分割するステップと、

データストリームデバイダユニット (D S D U) の特定の F I F O レジスタを介して各データ部分を通過させるステップと、

前記特定の F I F O レジスタからのデータ部分を、処理のためにグラフィックプロセッサユニット (G P U) 内の、前記特定の F I F O レジスタに割り当てられたストリーミングマルチプロセッサ (S M) に転送するステップと、

を備え、

前記データ部分は、前記 D S D U 内の、第 1 のアドバンスト・イクステンシブル・インタフェース (A X I) ユニットと、前記 G P U 内の、第 2 のアドバンスト・イクステンシブル・インタフェース (A X I) ユニットの介して転送される、

大容量のデータを効率的に処理する方法。

【請求項 5】

特定の F I F O レジスタから受信したデータ部分は、前記 D S D U 内の、前記 F I F O レジスタに割り当てられた第 1 の A X I ユニットと、前記 G P U 内の、前記 F I F O レジスタに割り当てられた第 2 の A X I ユニットの介して前記 G P U 内の、前記特定の F I F O レジスタに割り当てられた S M に転送される、請求項 4 に記載の方法。

【請求項 6】

F I F O レジスタから受信した前記データ部分の各々は、前記 D S D U 内の、共通の第 1 の A X I ユニットと、前記 G P U 内の、共通の第 2 の A X I ユニットの介して、前記 G P U 内の、前記特定の F I F O レジスタに割り当てられた S M に転送される、請求項 4 に記載の方法。

【発明の詳細な説明】

【技術分野】

【 0 0 0 1 】

膨大の量のグラフィックデータおよび計算を操作／処理するのに適したコンピューティングシステムは、典型的に、中央処理装置 (C P U) は別として、必要なデータを操作し、処理するために適合および指定された、G P U 、G P G P U 、D S P 、S I M D ベースプロセッシングユニット、V L I W ベース処理ユニットを備える。そのような計算システムの構造は、この分野においてよく知られている。この構造は、典型的に C P U と M P U との間で計算タスクを分割するので、重たい計算は、M P U に割り当てられ、残りの計算タスクは、C P U が操作するように残される。

【 0 0 0 2 】

しかしながら、よく知られた構造は、大量の図形データが含まれる場合、C P U と M P U との間で双方向に (back-and-forth) 図形データの転送を管理するのに必要な大量のハンドリングリソースにより、効率が低下する。いくつかの場合、C P U - M P U 計算構造におけるデータ計算に使用可能なネット時間は、5 % 未満の可能性はある。たとえば、N v i d i a (登録商標) コンピューティング・ユニファイド・デバイス・アーキテクチャ (Computing Unified Device Architecture) (C U D A) 並列コンピューティングプラットフォームと、アプリケーションプログラミングインターフェイスモデルの場合、グラフィックデータの処理に費やされる一般的な時間部分は、C P U 環境から G P U 環境 (C U D A メモリなど) にグラフィックデータを転送する場合 4 9 % 、グラフィックデータを G P U 環境から C P U 環境 (C U D A メモリ) に返送する場合、4 7 % であり、グラフィック計算の場合には、4 % 未満である。このような非常に低いグラフィック計算効率

10

20

30

40

50

は、グラフィックデータが、プロセッサ間で転送される方法を定義する共通の構造から来る。

【0003】

グラフィック計算に割り当てられた時間を、実質的に上昇させる、MPU効率の実質的な上昇を可能にする必要性がある。

【発明の概要】

【0004】

ストリーミングマルチプロセッサを介したグラフィックデータソースと、グラフィック処理ユニット（GPU）との間で交換されるグラフィックデータスループットを高める方法が開示される。GPUは、プロセッシングコアユニット（PCU）、レジスタファイルユニット、複数のキャッシュユニット、シェアドメモリユニット、統合キャッシュユニットおよびインタフェースキャッシュユニットを備えることができる。この方法は、インタフェースキャッシュユニットを介して、および複数のキャッシュユニットを介して、および統合されたキャッシュユニットを介して、グラフィックデータのストリームを、レジスタファイルユニットに転送するステップと、レジスタファイルユニットからのグラフィックデータの第2のストリームを、プロセッシングコアユニットに転送するステップと、レジスタファイルユニットを介してデータの頻繁に使用される部分をシェアドメモリに記憶し、受信するステップとを備えることができる。

【0005】

いくつかの実施形態において、レジスタファイルユニットは、PCUにより処理されたデータを、そのデータを使用する頻度のレベルに応じて、さらにデータを受信することができる限り、シェアドメモリに仕向けるように構成される。いくつかの実施形態において、使用する頻度のレベルは、PCUにより決定される。

【0006】

グラフィックデータを処理するように構成されるプロセッシングコアユニット（PCU）と、PCUからグラフィックデータを供給し、PCUからの処理したグラフィックデータを受信して一時的に記憶するように構成された、レジスタファイルユニットと、レジスタファイルユニットからグラフィックデータを供給し、前記レジスタファイルユニットからの処理されたグラフィカルデータを受信し、一時的に記憶するように構成された、複数のキャッシュユニットと、レジスタファイルユニットからグラフィックデータを供給し、レジスタファイルユニットから処理されたグラフィックデータを受信し一時的に記憶するように構成されたシェアドメモリユニットと、レジスタファイルユニットからグラフィックデータを供給し、レジスタファイルユニットからの処理されたグラフィックデータを受信し一時的に記憶するように構成された、統合キャッシュユニットと、そして、高速でグラフィック処理のためのグラフィックデータを受信し、共有メモリユニットと統合キャッシュユニットの少なくとも一方に、グラフィックデータを提供し、統合キャッシュユニットから処理済みグラフィックデータを受信し、処理されたグラフィックデータを外部処理ユニットに供給する。

【0007】

いくつかの実施形態では、グラフィックデータ要素の少なくともいくつかは、PCUによる近い呼び出し（close call）の確率に関連付けられた優先度に基づいて、共有メモリユニット内のPCUによる処理の前および／または後に格納される。いくつかの実施形態において、優先度は確率が高くなるにつれ高くなる。

【0008】

データストリーム分割ユニット（DSDU）およびグラフィックス処理ユニット（GPU）を備えた、未処理データを処理する回路が開示される。DSDUは、データのストリームを受信し、データの一部に分割し、データの一部の各々を、複数の先入れ先出し（FIFO）レジスタの1つを介して通過させるように構成された、複数のFIFOレジスタと、前記データ部分を受信するように構成された、アドバンス・トイクステンシブル・インタフェース（Advanced Extensible Interface）（AXI）とを備える。GPUは、第

1のAXIユニットからデータ部分を受信するように構成された、第2のアドバンスト・イクステンシブル・インタフェース（AXI）ユニットと、それぞれのFIFOレジスタから各データ部分を受信するように構成され、受信したデータ部分进行处理する複数のストリーミングマルチプロセッサ（SM）を備える。

【0009】

いくつかの実施形態において、DSDU内の特定のFIFOレジスタは、DSDU内に割り当てられた第1のAXIユニットを介して、およびGPU内に割り当てられた第2のAXIユニットを介して、GPU内に割り当てられたSMに接続される。いくつかの実施形態において、DSDU内のFIFOレジスタの各々は、DSDU内の第1のAXIユニット、およびGPU内の共通AXIユニットを介してGPU内に割り当てられたSMに接続される。

10

【0010】

未処理データのストリームを受信するステップと、前記ストリームを複数のデータ部分に分割するステップと、各データ部分をデータストリームデバイダユニット（DSDU）内の特定のFIFOレジスタを介して通過させるステップと、前記特定のFIFOレジスタからの前記データ部分を、処理のために、グラフィクスプロセッサユニット（GPU）内に割り当てられたストリーミングマルチプロセッサ（SM）に転送するステップと、を備えた、大量のデータを効率よく処理するための方法が開示される。いくつかの実施形態において、データ部分は、DSDU内の第1の特定のアドバンストイクステンシブルインタフェース（AXI）と、GPU内の第2の特定のアドバンストイクステンシブルインタフェース（AXI）を介して転送される。

20

【0011】

いくつかの実施形態において、特定のFIFOレジスタから受信したデータ部分は、DSDU内に割り当てられた第1のAXIユニットと、GPU内に割り当てられた第2のAXIユニットに転送される。いくつかの実施形態において、DSDU内のFIFOレジスタから受信したデータ部分の各々は、DSDU内の共通の第1のAXIユニット、およびGPU内の共通の第2のAXIユニットを介して、GPU内に割り当てられたSMに転送される。この発明としてみなされる主題は、明細書の結論部分で特に指摘され、明確に特許請求の範囲に記載される。しかしながら、この発明は、動作の組織および方法の両方に関して、それらのオブジェクト、特徴および利点と一緒に、添付した図面と共に読むとき以下の詳細な記載によって最もよく理解することができる。

30

【図面の簡単な説明】

【0012】

【図1】図1は、GPUを用いたコンピューティングユニットのデータフローを概略的に説明する。

【図2】図2は、GPUユニット内の典型的なストリーミングマルチプロセッサ（SM）の概略ブロック図である。

【図3A】図3Aは、この発明の実施形態に従って構成され、動作可能な未処理データ（UPD）処理ユニット（UPDHU）300を描画する概略ブロック図である。

【図3B】図3Bは、図3AのUPDHU300のような、UPDHUの一実施形態の概略ブロック図である。

40

【図3C】図3Cは、図3AのUPDHU300のような、UPDHUの他の実施形態の概略ブロック図である。説明の簡便かつ明瞭さのために、図に示されるエレメントは必ずしも縮尺通りではない。例えば、エレメントのいくつかの寸法は、明瞭さのためにエレメントに対して誇張されている可能性がある。さらに、適切に考慮する場合、参照符号は、図面間で対応するまたは類似するエレメントを示すために反復可能である。

【発明を実施するための形態】

【0013】

以下の詳細な記載において、この発明の完全な理解を提供するために、多くの特定の詳細が述べられる。しかしながら、当業者には、この発明は、これらの特定の詳細なしに実

50

施可能であることが理解されるであろう。他のインスタンスにおいて、よく知られた方法、手続およびコンポーネントは、この発明を不明瞭にしないように詳細に記載していない。既知のコンピューティングシステムにおけるCPU-GPU相互作用のボトルネックは、大部分が、CPUによってグラフィック関連データをGPUに仕向け、および処理したグラフィックデータをGPUから受信するのに使用される、データ転送チャンネルにある。典型的に、CPUとGPUプロセッサは、標準のコンピューティング環境で動作し通信する。

【0014】

図1を参照すると、GPUを用いてコンピューティングユニット100のデータフローを概略的に説明する。コンピューティングユニット100は、CPU111、CPUダイナミックRAM (DRAM) 111A、(メインボードチップセット) のようなコンピューティングユニット周辺制御ユニット112を備える。ユニット100はさらに、ユニット112を介してCPUとデータを通信するGPUユニット150を備える。GPUユニット150は、典型的に、ユニット112とGPUプロセッサとの間でデータをインタフェースするGPU DRAMユニット154、GPU処理ユニットのためのデータをキャッシュするように適合したGPUキャッシュユニット156 (例えば、L2 キャッシュユニット)、およびGPU処理ユニット158 (例えば、ストリーミングマルチプロセッサ/SM) を備える。

【0015】

処理ユニット100に入力され、GPU150により処理されるように意図されるグラフィックデータのフローは、データフロー (DF) の矢印により記載される。第1のデータフローDF1は、コンピューティングユニット100へのデータのフローを描画し、CPU111は、DF2のフローを、周辺制御ユニット (PCU) 112を介してDRAM 111Aに仕向け、そこから戻って、DF3-PCU112-DF4を介してGPU150に仕向ける。GPU150において、データは、DRAMユニット154を通り、キャッシュユニット156を通して、複数のストリーミングマルチプロセッサ (SMs) ユニット158に到達し、そこで、グラフィック処理が行われる。出来るだけデータフローボトルネックを無くすことは、この発明に従う方法と構造のターゲットである。

【0016】

図2を参照すると、GPUユニット内の典型的ストリーミングマルチプロセッサの概略ブロック図である。SM200は、プロセッシングコアユニット210 (コンピュータ・ユニファイド・デバイス・アーキテクチャ (CUDA) コアと呼ばれるときもある) と、コア201とキャッシュユニット230 (コンスタントキャッシュ)、250 (ユニファイドキャッシュ)、およびシェアドメモリ240を備える。SM200に入ってくるデータおよびそこから出ていくデータは、図1のGPUキャッシュユニット256 (例えば、キャッシュユニット156 (L2)) を用いて交換される。グラフィック処理が既知の方法で実行されると、GPUユニットは、処理されるデータ量全体が、グラフィック処理が開始される前にいくつかのSM200ユニットのメモリユニットにロードされるまで待つであろう。

【0017】

データ転送時間を低減する1つの方法は、データ転送を最小に低減することである。例えば、コア210により計算された中間結果は、DRAMに記憶する代わりにレジスタファイル220に記憶することができる。さらに、シェアドメモリ240は、通常行われるようにアウトバウンド (outbound) で循環させる代わりに、SM200内で頻繁に使用されるデータを記憶するために使用することができる。いくつかの実施形態において、使用頻度のレベルは、PCUにより決定される。さらに、コンスタントメモリユニットおよび/またはキャッシュメモリユニットは、SM210で定義することができる。この発明のさらなる実施形態によれば、CPUコンピューティング環境とGPUコンピューティングとの間のデータフローボトルネックは、CPUを、グラフィック関連データをすべて処理するように特に構成されたコンピューティングユニットと置き換えることにより、低減

することができるか、または消去することができる。

【0018】

図3Aを参照すると、この発明の実施形態に従って、構成され、動作可能な未処理データ(UPD)ハンドリングユニット(UPDHU)300を描画する概略ブロック図であり、図3Bおよび図3Cは、図3AのUPDHU300のようなUPDHUの2つの異なる実施形態350および380の概略ブロック図である。ここで使用される「未処理データ」という用語は、処理しようとしているデータの大きなストリームに関連し、典型的には、大きな計算容量を必要とし、例えば、仮想的に「リアルタイム」で(すなわち、できるだけ小さい待ち時間で)処理する必要がある、グラフィックデータの高速ストリーム(例えば、4Kビデオカメラから受信された)である。図3Aに描画されるUPDHU300のアーキテクチャは、CPU-GPUの既知のアーキテクチャに典型的な、データストリームの固有のボトルネックを克服するように設計され、データ獲得の入力ストリームは、最初に、CPUにより処理され、次に、一時的にCPUに関連付けられたCPUメモリ、および/またはRAMに記憶され、次に、(例えば、周辺コンポーネント相互接続イクスプレス(PCIe)バスを介して)GPUに再び転送され、GPUの一部である複数のストリーミングプロセッサに送信される前に、GPUプロセッサにより再び処理される。

10

【0019】

図3Aに関して、ここで記載された例は、アドバンスド・イクステンシブル・インタフェース(AXI)に従って、動作するようにプログラムされたフィールドプログラマブルゲートアレイ(FPGA)の使用を示しているが、当業者には、ここに記載した動作の方法は、それぞれのGPUとインタフェースするように適合し、および高いスループットで、グラフィック関連データを大量に転送するのに適した、他のコンピューティングユニットを用いて具現化することができる。

20

【0020】

この発明の実施形態によれば、データストリーム・デバイダ・ユニット(DSDU)304は、例えば、大量のストリーミングUPD、例えばカメラからのビデオストリームを受信し、それを複数のより小さなストリームに分散し、GPUのSMsに転送するようにプログラムされた、FPGAを用いて具現化することができる。FPGAとGPUは、さらに、GPUの複数のSMsの、少なくとも1つのSMが完全にロードされるやいなや、GPUが、転送されたグラフィックデータの処理を開始するように動作するようにプログラムすることができる。ほとんどの場合、完全にロードされたSMsは、フルデータファイルよりも小さいデータ量を保持し、それゆえ、GPUによる処理は、この実施形態に従って開始され、データファイル全体がGPUにロードされた後で、処理を開始する一般的に知られている実施形態に比べて、はるかに速い。

30

【0021】

一例示実施形態において、UPDHU300は、複数の先入れ先出し(FIFO)レジスタ/ストレージユニットアレイ304A(FIFOユニットは、個別に示されていない)を備えたDSDU304を備え、そのうちの1つのFIFOユニットを、GPU320のSMs318の各々に割当てることができる。いくつかの実施形態において、DSDU304により受信されるUPDストリームは、複数のデータユニットに分割され、FIFOユニット304Aを介して、GPU320に転送することができ、AXIインタフェースのようなインタフェースユニットを介して、GPUにブロードキャストすることができ、それにより、各FIFO304Aのデータユニットが、関連するSM318に転送され、それにより、例えば、シングルアクション・マルチデータ(SEVID)コンピューティングを可能にする。GPU320の各(単一でも)SM318は、AXIインタフェースを介して、関連するFIFO304Aから受信した未処理データのそれぞれの部分がロードされるとき、GPU320は、処理を開始することができ、UPDファイル全体がロードされるまで待つ必要はない。

40

【0022】

MSU310は、長いストリームのグラフィックデータを受信するように構成された、

50

未処理データインタフェースユニット302を備えることができる。インタフェースユニット302を介して受信した大量の未処理データは、より小さなサイズの複数の数のデータユニットに分割され、FIFOユニット304A内の割り当てられたFIFOユニットを介して転送され、次に、AXIチャネル315で、GPU AXIインタフェースを介して、GPU320の割り当てられたSM318に転送される。SMs318のそれぞれのSMにより処理されたデータユニットは、次に、AXI接続を介して、MSUに転送することができる。上記のように、上述した実施形態では、CPU-GPUアーキテクチャにおいて、典型的に大きなオーバヘッドが節約される。

【0023】

図3Bおよび3Cは、この発明の実施形態に従う、図3AのMSU310を具現化する2つのオプションのアーキテクチャの概略ブロック図を描画する。図3Bは、MSU350がFIU356とGPU358を備えることを描画する。FIU356は、複数のFIFOユニット（集合的に356Aと名前が付けられている）-FIFO0、FIFOi・・・FIFO_nを備える。各FIFOユニットは、割り当てられたFPGA AXI（F-AXI）ユニット-FAXI0、FAXIi・・・FAXIn（集合的に356Bと名前が付けられている）とアクティブ通信中であり得る。別個のF-AXIユニットの各々は、間接的に、割り当てられたGPU AXI（G-AXI）ユニット-GAXI0、GAXIi・・・GAXInとダイレクトに接続し得る。G-AXIインタフェースユニットの各々は、割り当てられたSM-SM0、SMi・・・SM_nとアクティブに接続することができ、データを供給することができる。

【0024】

さらに他の実施形態によれば、図3Cに示すように、MSU380は、FIU386とGPU388を備える。FIU386は、複数のFIFOユニット（集合的に386Aの名前がつけられている）-FIFO0、FIFOi・・・FIFO_nを備えることができる。各FIFOユニットは、複数のFIFOユニットから単一のAXIストリームへのデータストリームの管理を制御するように構成することが出来るFPGA AXI（F-AXI）ユニットとアクティブに通信することができる。AXIストリームは、GPU388のAXIインタフェースに送信して、次にそれぞれのSMsユニット-SM0、SMi、・・・SM_nに分割することができる。図3Bに描画されるアーキテクチャは、より高速な全体性能を提供することができるが、（記載した回路を具現化する集積回路（IC）のために）より多くのピンと、より多くのワイヤ／導管（conduits）を必要とする場合がある。図3Cに描画される、アーキテクチャは、相対的により遅い全体性能を提供するが、（記載した回路を具現化する集積回路（IC）のために）より少ないピンと、より少ないワイヤ／導管（conduits）しか必要としない。

【0025】

上述した、デバイス、構造、および方法は、既知のアーキテクチャおよび方法に比べて、大量の未処理データの処理を加速することができる。例えば、既知の実施形態において、処理／アルゴリズムがHPU上で開始できる前に、画像全体を転送する必要がある。画像サイズが1GBの場合、GPUへデータを転送するPCI-Eバスの理論上のスループットは、32GB/sであり、待ち時間は、 $1\text{GB} / (32\text{GB/s}) = 1/32\text{s} = 31.125\text{ms}$ 31.3msである。それに反して、この発明の実施形態に従うFPGAを用いると、すべてのSMユニットをフルにロードする必要があるだけである。たとえば、Tesla P100 GPUには、56個のSMユニットがあり、各SMには、32ビット（単精度モード）をサポートする64コアまたは64ビット（拡張精度モード）をサポートする32コアがあるため、完全にロードされたGPUのデータサイズ（単精度モードまたは拡張精度モードで同じ結果）は、 $56 * 32 * 64 = 114688\text{ビット} = 14.336\text{MB}$ バイトである。FPGAからGPU AXIストリームへの理論的スループットは、896MB/s（56レーンの場合）であり、待ち時間は、 $14.336\text{MB} / (896\text{MB/s}) = 14.336 / 896\text{s} = 16\text{ms}$ であり、これは実質半分の待ち時間である。

10

20

30

40

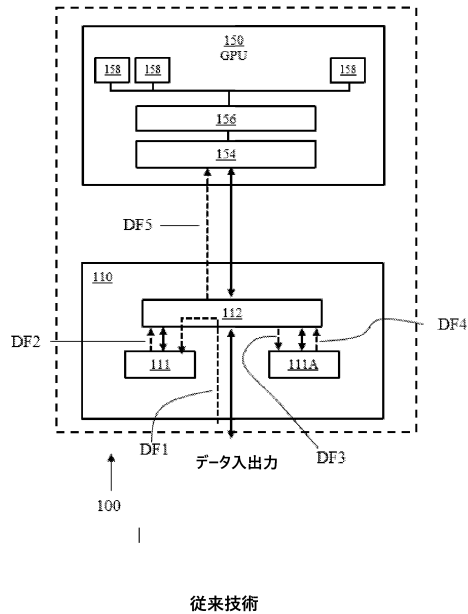
50

【 0 0 2 6 】

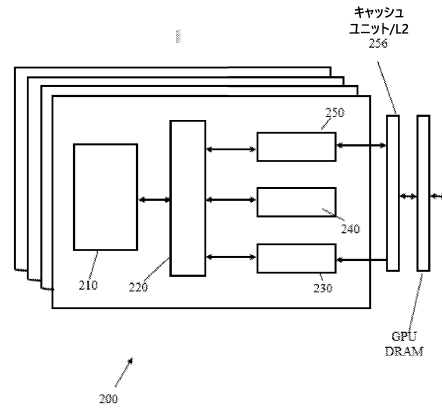
本明細書では、本発明の特定の特徴を例示および説明してきたが、多くの修正、置換、変更、および同等物が当業者に思い浮かぶであろう。したがって、添付の特許請求の範囲は、本発明の真の精神の範囲内に、あるそのようなすべての修正および変更を網羅することを意図していることを理解されたい。

【図面】

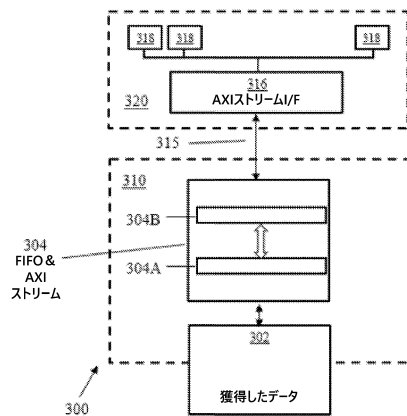
【図 1】



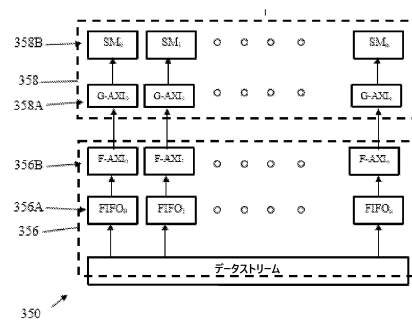
【図 2】



【図 3 A】



【図 3 B】



10

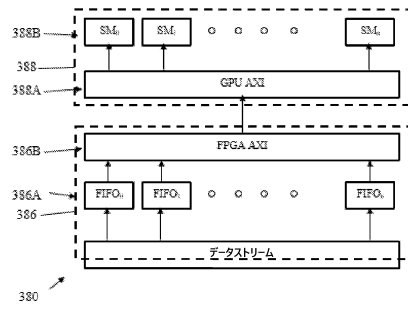
20

30

40

50

【図 3 C】



10

20

30

40

50

フロントページの続き

弁理士 小菅 一弘
(72)発明者 ハニア シャハール
イスラエル国 4485600 ケドゥミム アリエル ストリート 7
(72)発明者 ゼルツァー ハナン
イスラエル国 4526514 ホド ハシャロン ブネイ ブリット ストリート 12
審査官 北村 学
(56)参考文献 特表2003-502958 (JP, A)
特開2012-146323 (JP, A)
特開平06-324997 (JP, A)
特開2006-091982 (JP, A)
特開2012-150735 (JP, A)
特開2015-156645 (JP, A)
(58)調査した分野 (Int.Cl., DB名)
IPC G06F 12/08 - 12/128
G06F 9/38
G06F 13/38 - 13/42