

AI, 脳シミュレーションと計算機

牧野淳一郎

理化学研究所 計算科学研究機構

エクサスケールコンピューティング開発プロジェクト

コデザイン推進チーム チームリーダー

はじめに

すみません、今日の話はほぼ 11/28 の理研シンポジウムの使い回しです。

牧野の話のポイント

- 何故 AI が話題になるか？
- 基本的にはやはり計算機的能力(計算速度、記憶能力)の発展によって可能になってきたから。
- なので、ここでは、計算機的能力の発展はどう進んできたか、今後どうなるか、どうすべきか、それと脳シミュレーション・人工知能研究の関係はどうか、なにを進めるべきか、という話をしたい。
- 私は基本的に保守的なので、CMOS デバイスが当面主流、ということが前提の話をする。
- それでもまだ数桁の能力発展の余地があるので。

今日の話は、計算科学研究機構エク
サスケールコンピューティング開発
プロジェクトの見解を反映したもの
ではありません

話の構成

- これまでの計算機の進歩と現状の困難
- 解決の方向
- 「人工知能」に必要な計算機は？
- 深層学習と行列乗算
- 可能なアーキテクチャ
- まとめ

これまでの計算機の進歩

1940年代から現在までの70年間、ほぼ10年で100倍。何故そのような指数関数的進歩を長期に続けたのか?(今後はどうか?)
基本的な理由:

- 使うスイッチ素子が高速になった
- 使うスイッチ素子が小型、低消費電力になって、沢山使えるようになった
- 使うスイッチ素子が安くなって、沢山使えるようになった
(スパコンの物理的大きさは70年代が最小。そこまで段々小さくなって、そこからまた大きくなった)



素子の高速化

- といっても、真空管でもそれなりに速かった。
- スイッチング速度が重要でないわけではないが、配線を信号が伝搬する速度のほうが昔から重要。
- 昔は信号はほぼ光の速さでつたわった。
- 最近の LSI 上の配線は非常に細く (抵抗が大きく)、キャパシタンスを充電しないといけないことによる RC 遅延のため、信号が伝わる速度は光速度よりはるかに低い。
- 太い配線に大電流を流せば速いが、大量の電力消費になる

素子の小型化

- 真空管 → トランジスタ → IC という進化は 1970 年代までは重要
- サイズだけでなく、消費電力が下がることが重要
- 80 年代から重要になったのは CMOS LSI の微細化。10 年でサイズが 1/10 になる
- CMOS 素子では (2000 年くらいまでは) 微細化すると電圧を下げることができ、消費電力が下がり、速度は向上した。いわゆる CMOS スケーリング。
- 2000 年頃からは電圧が下がらないので、電力はちょっと下がるが速度は上がらなくなった。いわゆる CMOS スケーリングの終焉。
- そろそろ微細化も困難になってきた。また、トランジスタの構造・製造工程が複雑になり、微細化するとかえって価格上昇するようになった。いわゆるムーアの法則の終焉。

素子の性能向上、サイズ低下と スパコンの性能向上

「スパコンの進歩」は4つの時期に分けられる

I スパコンが完全パイプライン化した乗算器をもたない時期
(1969年まで)

II スパコンは1つ以上の演算器をもつが、1チップにはまだ
演算器1つが入らない時代 (CMOS では1989年まで)

III 1チップに複数の演算器がはいるが、微細化すると動作ク
ロックが上がり、電力が下がった時代 (2001年頃まで)

IV チップに多数の演算器がはいるが、微細化してもクロック
が上がり、電力がへらない時代 (2001-)

このそれぞれで、素子の性能向上がどう使われたかは違う。

I期: 1 演算器未満の時代 ～ 1969

- この時期には、メモリアクセスのほうが演算より速い
- 演算器の性能向上が最重要課題
- CDC 6600 くらいまで。
- 計算機の構成方法もまだ手探り。
- CDC 7600 で1 演算器はいるようになる

II 期: 1 演算器以上の時代 ～ 1990

- 1つ以上の演算器を有効に使うことが課題
- パイプライン化 (ベクトル命令)、複数の演算ユニットの SIMD and/or MIMD 並列実行が必要になる
- 演算器の数が増えると、メモリとどうつなぐかが課題になる。
- 演算器 16-32 個で破綻する (1992 年頃): メモリと演算器の間のデータ移動回路が大規模・複雑になり過ぎるため
- 末期の計算機: Cray T-90, 日立 S-3800
- 富士通 VPP500 では、分散メモリにすることで当座しのごうとした:ある程度の成功
- 他の方向: 1チップマイクロプロセッサでの分散メモリ。プロセッサ間の通信は細い線でいいことにする。これが90年代以降の主流になった

III期: 1チップ高速化の時代 ～ 2005

- 1チップマイクロプロセッサを分散メモリで多数つなぐの
は II期末期から
- 1チップに1演算器以上がはいるが、ありあまるトランジ
スタを演算器を増やすのにはつかわないで動作クロックの
向上、キャッシュメモリの大型化に使った時代
- 1990年代。 牧野の意見としては「失われた10年」
- 計算機全体としては II期に起こった問題がチップレベル
でも起こるのを先送りにした
- メモリとの接続は速度が不足になった (memory wall): キ
ャッシュでごまかす方針
- 迷走したプロセッサ開発も多い:各社の複数チップ共有メ
モリプロセッサ。(失敗プロジェクト: ASCI Red 以外の
ASCI マシン)
- 末期の計算機: Intel Pentium 4, DEC Alpha 21264

IV 期: 1チップ多演算器化の時代 ～ 2020?

- 引き続き、1チップマイクロプロセッサ、分散メモリ。
- デザインルール 130 nm あたりから、動作電圧低下、速度向上に限界
- マルチコア化、コア内 SIMD 化を同時に進めている
- 80年代のベクトルスパコンの辿った道を追いかけている
- ということは、16-32コアで破綻がくるはず
- 破綻がくることの予見: Intel Xeon Phi (60コア)
- GPGPU はちょっと別。

まとめると

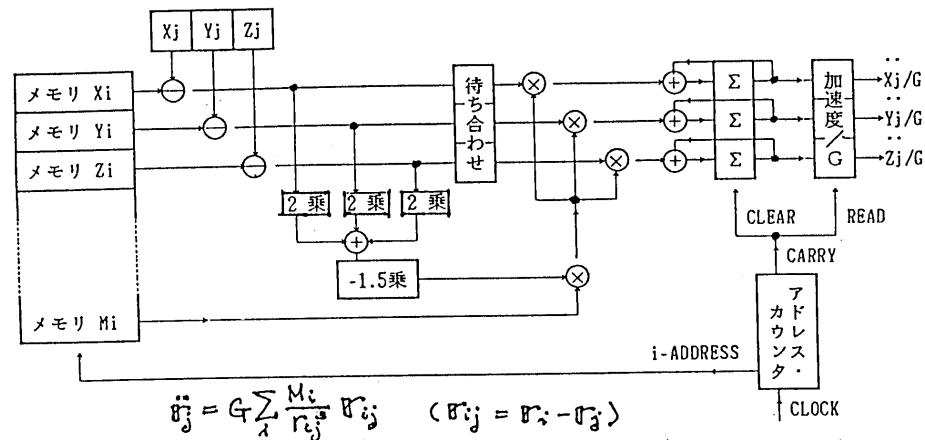
- 半導体技術は 2000 年頃に CMOS スケーリングが終焉、さらに微細化自体のスローダウンが進行中であり、もはや微細化が技術進歩を牽引していない。
- 昔の Cray-1 みたいな「スパコン」の進歩は、プロセッサコア 32 個くらいの並列度で限界、破綻した
- 現在のマイクロプロセッサは、それくらいのコア数に達してきており、破綻が近い
- この 2 重の困難をどう解決するか、という展望が必要。

「解決」の方向

- 90年代にベクトル並列機が向かった方向を追いかける。
 - VPP500 を1チップ化。チップ内でも物理共有メモリを断念して、分散メモリにする
 - GPU は(外部に共有メモリはあるが)これに近い。
 - もっと極端なアーキテクチャが望ましい(GRAPE-DR、ポスト「京」向け「加速部」)
 - これはこれで人工知能応用もできるが今日はいあまりこの話はしない方向で、、、
- 専用回路にいく
 - 半導体が進歩しなくなるので専用回路が陳腐化しない。
 - 同じテクノロジーで、アプリケーションによるが GPU の 5-50 倍程度は電力性能あげられる。
 - 汎用プロセッサに比べると開発が非常に容易(「加速部」に比べても楽)。

専用回路の例

- GRAPE, MD-GRAPE: それぞれ重力多体問題、分子動力学の専用計算機。粒子間相互作用を専用パイプラインで計算するのが基本的発想。汎用プロセッサとI/Oバスでつないで計算分担 (MDGRAPE-4 以前は)
- ANTON: 専用パイプライン+汎用プロセッサ+チップ間通信ネットワークを1チップに集積、通信レイテンシを減らすことで比較小さい系に対して超高速計算を実現。

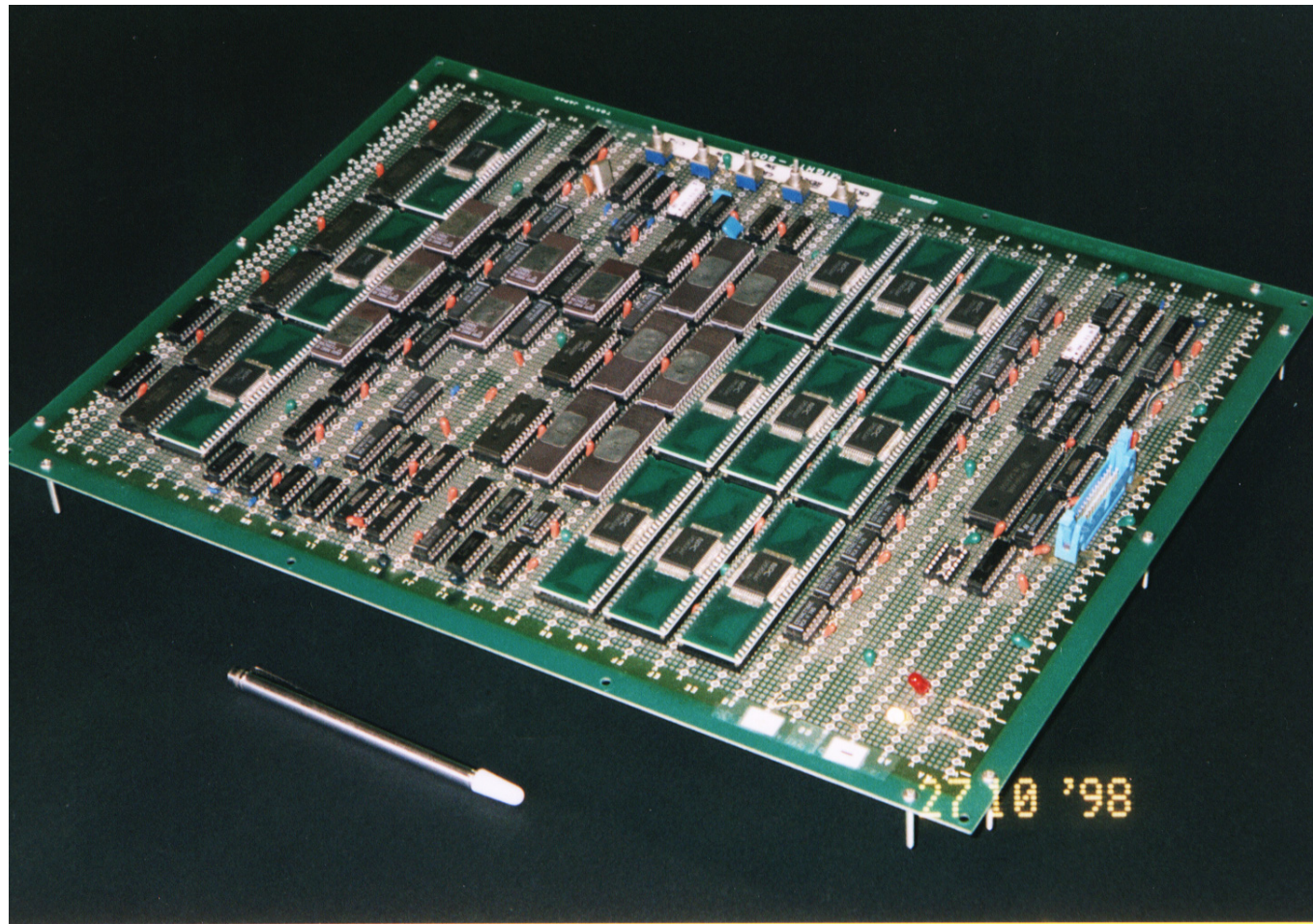


+, -, ×, 2乗は1 operation, -1.5乗は多項式近似でやるとして10operation 位に相当する。
 総計24operation.
 各operation の後にはレジスタがあって、全体がpipelineになっているものとする。
 「待ち合わせ」は2乗してMと掛け算する間の時間ズレを補正するためのFIFO (First-In First-Out memory).
 「Σ」は足し込み用のレジスタ。N回足した後結果を右のレジスタに転送する。

図2. N体問題のj-体に働く重力加速度を計算する回路の概念図。

近田、1988

GRAPE-1(1989)



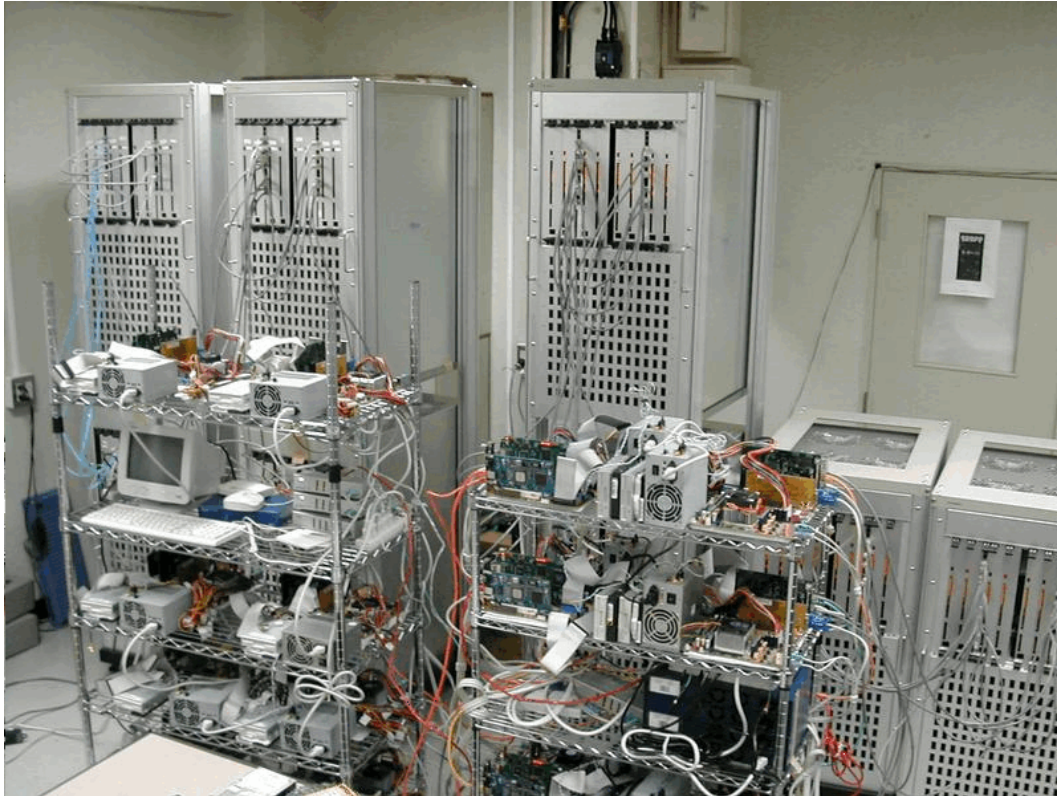
GRAPE-4(1995)



1Tflops を世界に
先駆けて実現

汎用で 1TF 超えた
ASCI Red の2年
前、1/20 のコスト

GRAPE-6(2002)



2002年の 64 Tflops
システム

初代地球シミュレー
タの 1.5 倍の性能
を 1/100 のコスト
で

そうはいっても GRAPE で 人工知能も脳もできないのでは？

- それはそうである。
- では人工知能向けの専用回路って？
- そもそも最近の人工知能研究ってどんな計算しているの？
- 脳シミュレーションは？

人工知能のキー概念：深層学習

深層学習における計算

- 深層学習は多段ニューラルネット。
- ニューラルネット：入力に係数掛けて足したものが次段のニューロンの入力。係数は同じ段でニューロン毎に違う = 行列ベクトル積。入力が複数あれば行列積。
なので：
 - 計算の 99.9% くらいが行列乗算 (BP 学習でも)
 - GPU が使われてるのは結局単精度行列乗算専用マシンとして
 - NVIDIA が倍精度なしの Maxwell も高く売るのが始めたくらいには需要がある。
- 学習は多段ニューラルネットに対して BP が最適かどうかは (私には) 良くわからない。まあでも多分行列乗算ができればいいような方法が今後も使われる。

つまり：深層学習専用回路＝行列乗算専用回路

行列乗算専用回路

- 意外に重要なアプリケーションが色々ある
 - 量子化学計算一般
 - 深層学習
 - FEM (DDM とか使うと結局、、、)
 - stiff な系が多数ある話 (東大の心臓シミュレーションとか)
- 28nm で、倍精度乗算器+加算器だけなら 0.02 平方ミリくらい。300 平方ミリで 1 万ペア。500MHz で回して 10TF、100-150GF/W くらいまでいけるはず。単精度なら 40TF、300-500GF/W。
- 28nm の GPU の 10 倍くらいの電力性能。
- 多分 2025 年くらいまで使える。
- (もちろん Green 500 も狙える)

Green 500 #1 マシンは理研にある

広報活動

[Home](#) > [広報活動](#) > [トピックス2015](#) >

トピックス

[← 前の記事](#)

2015年8月4日

理化学研究所

株式会社ExaScaler

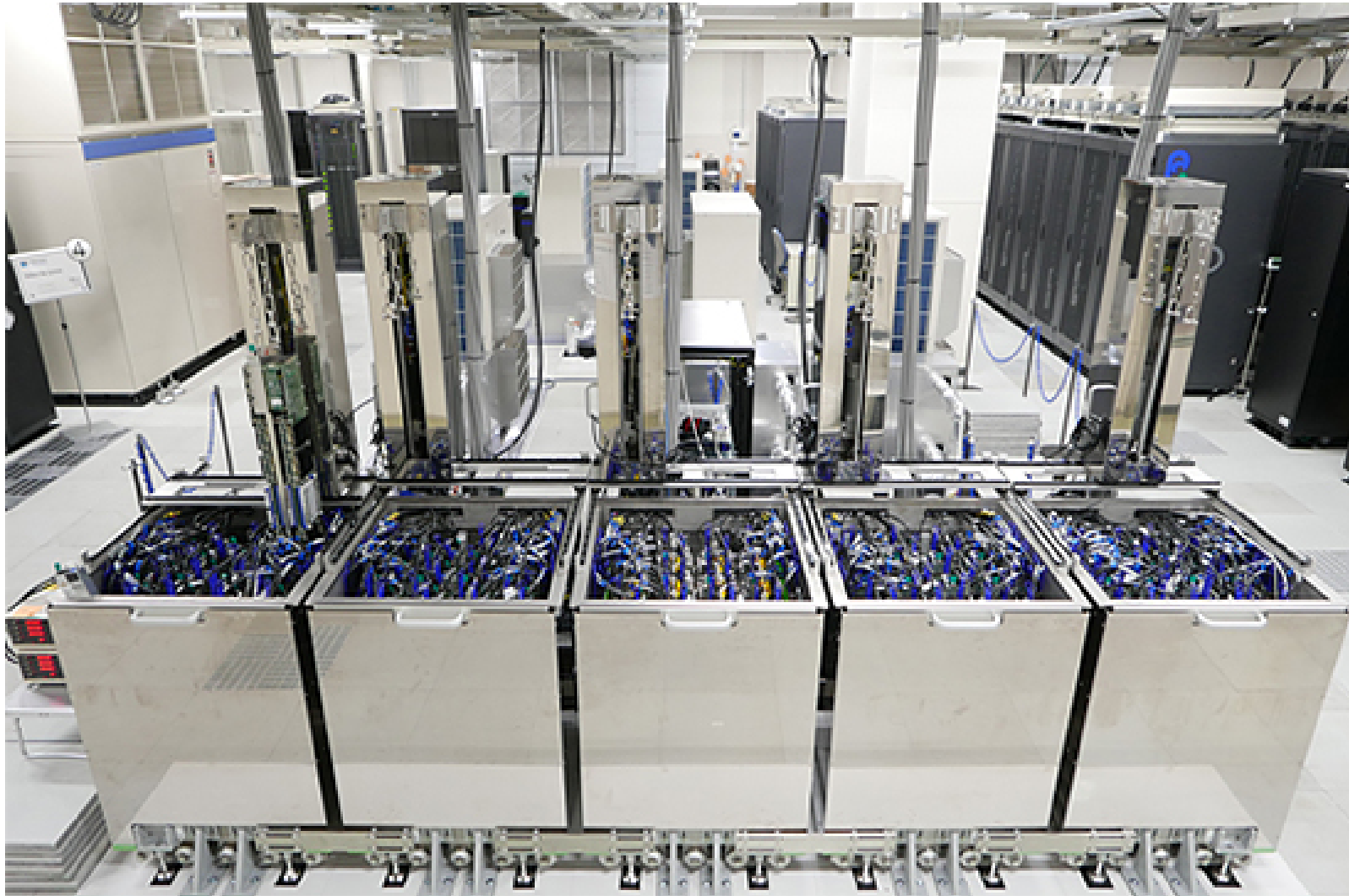
株式会社PEZY Computing

 いいね！

スーパーコンピュータ「Shoubu（菖蒲）」がGreen500で世界第1位を獲得

理化学研究所（理研）が、株式会社ExaScaler、株式会社PEZY Computingと共同で一部設置を完了させた2 PetaFLOPSi

PEZY/Exascaler “Shoubu” システム



深層学習専用にもうちょっと 頑張ってみると、、、

数千 × 数千 の行列積を「可能な限り短い時間で」やりたい。

- 大量データに対する BP 学習はあんまり上手く並列化できてない。
- 1つのニューラルネットを多数の GPU でとかは全然無理
- チップ間を高速ネットワークでつなげばなんとかならないか？

例

- 1チップ 30TF (単精度)
- 行列サイズ 2048
- 36 チップを 6×6 の2次元グリッドに
- 行列乗算が 16 マイクロ秒で終わる。
- 1方向の転送速度が 16MB を 16 マイクロ秒で、1TB/s。
6でわって170GB/s
- すぐ隣のチップとならできるかも。5GHz x 300本。
- GPU より 100 倍速い。

チップ間通信

- 170GB/s はまあ今どきなら、、、 25Gbps SERDES なら 50 ペアくらい。
- 隣のチップくらいまでなら embedded clock 使わないパラレル転送でも 10Gbps はできる。
- もちろん PEZY ではもっと高密度な磁界結合を開発中。
- データ圧縮は多分可能。音声や動画データでは時間相関があるので、予測演算が可能なはず。
- DNN の中間層でそんなことができるかどうかは自明ではないが、、、

コストの話

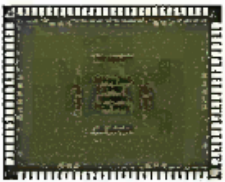
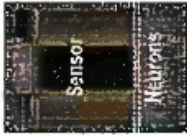
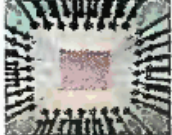
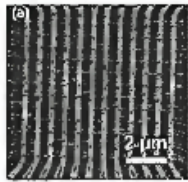
- LSI 開発の問題点は初期コストが高騰したこと。現在、28nm でマスク代だけで1億はだいぶ超えるはず。
- 設計を渡すインターフェースを RTL として、高い IP あまり使わないとして 4-5 億くらい。
- 製造コストはチップあたり 20-30 万。行列乗算回路の他に色々いれるのもうちょっと上がるかも。倍にはならない(ように設計する)。

ニューロン専用チップは？

日経エレクトロニクス 2016/3

表2 各種の抵抗変化素子を用いてシナプス可塑性を実現した脳型チップの例

(写真：パナソニックはVLSI Symposium 2013、それ以外はIEDM 2015)

開発元	IBM社	パナソニック	POSTECH、SK Hynix社など		UCSB、デンソーなど
発表時期	2015年12月* ¹	2013年6月* ²	2015年12月* ¹		2015年5月、同12月* ¹
チップ、または回路の写真			ニューロン 	シナプス 	
実装上の特徴	シナプスを2T-1R PCMで構成し、発火と学習のモードを分けた。プロセスは90nm世代	CMOS型ニューロンの配線層にシナプス (FeMEM) を形成	ニューロンとシナプスを別々のチップに実装		CMOSニューロンの上に12×12クロスバー構造のシナプスを形成
ニューロン数	256個	9個	1400個以上* ³		144個
シナプス数	6万5536 (256×256) 個	144個	8K個、または11K個* ⁴		144個
シナプス可塑性の実現素子	GeSbTeを用いたPCM	FeMEM	TiN/PCMOメモリスター		Al ₂ O ₃ /TiO _{2-x} メモリスター
ネットワークのタイプ	不定	Hopfieldネットワーク	パーセプトロン		多層パーセプトロンやCNN

FeMEM : Ferroelectric MEMristor PCM : 相変化メモリー PCMO : Pr_{0.7}Ca_{0.3}MnO₃

*1 IEDM 2015 *2 VLSI Symposium 2013 *3 シナプスとは別チップに実装し、変更可能 *4 Kは1024

ニューロン専用チップは？

- 実際問題としてニューロン数もシナプス数も少ない。使えるプロセスが古い。速度も？
- 保守的な牧野としては、CMOS デジタル回路はこの辺でも主流にとどまるのではないかと、、、

CNN だけでいいのか？

CNN だけでいいのか？

どうなのでしょう？

まとめ

- 半導体の進歩がスローダウン、終焉を迎えつつあり、また汎用マイクロプロセッサのアーキテクチャの進歩も、90年代に共有メモリベクトル並列プロセッサが辿った道と同じところで限界にきている。
- 今後の計算機アーキテクチャはとにかく電力性能が第一。GPU/加速部的アプローチか専用回路アプローチ。
- 深層学習ならおそらく圧倒的に専用回路(行列乗算専用回路)が有利。消費電力 1/10 以下、1 ケースの学習を現状の100倍高速化することも可能。GPU では将来になってもあんまり速くなりそうにない。
- チップ開発の費用は安くはないが、先端プロセスではないので比較すれば安価。
- アナログニューロンチップは将来性として難しいのでは、、、