

GRAPHCORE IPU (INTELLIGENCE PROCESSING UNIT) IPU-POD: NEW BREAKTHROUGHS IN AI AT SCALE

AI専用 革新プロセッサIPU

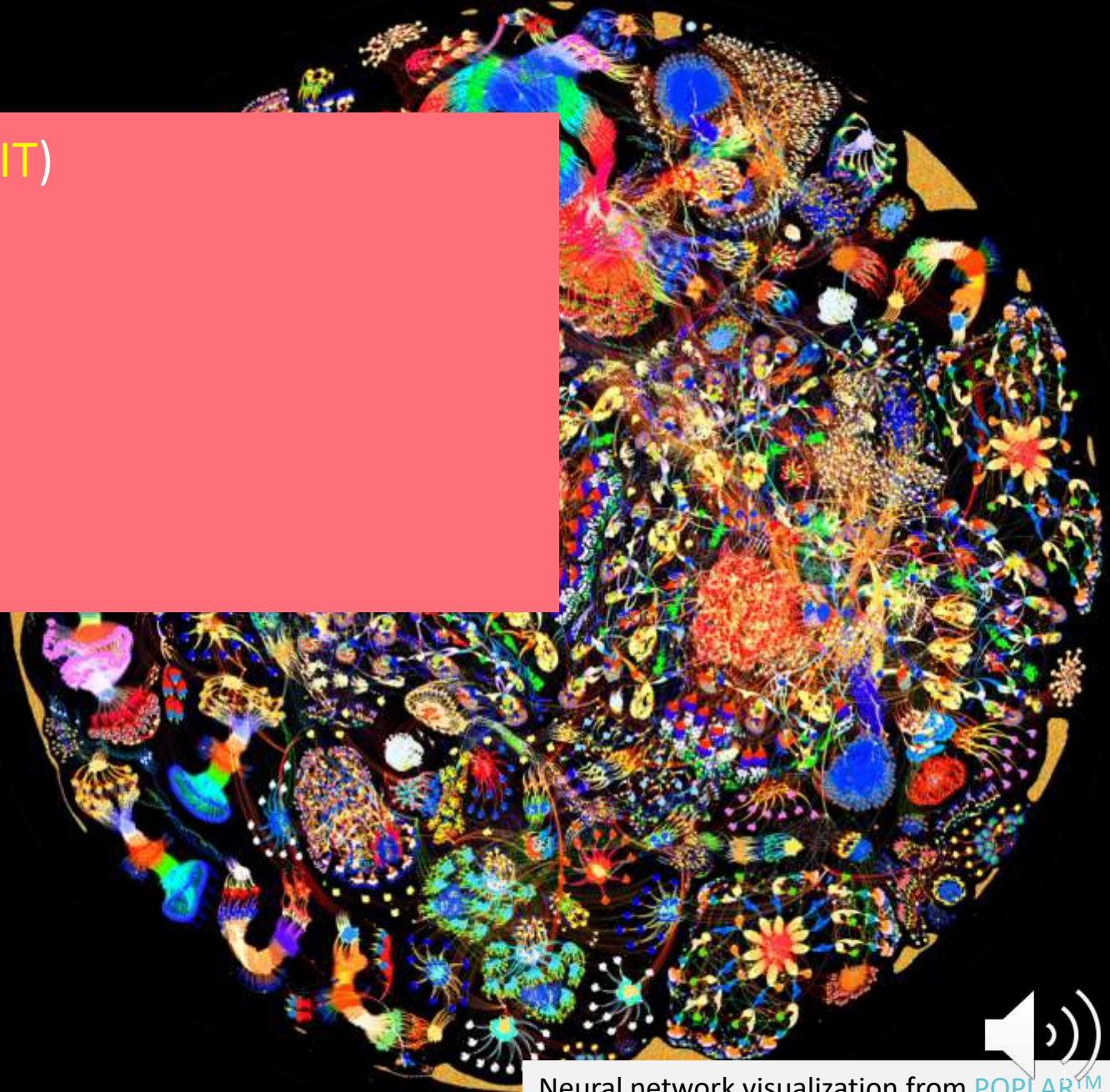
中野 守 mamorun@graphcore.ai

Jan 20 2022

Graphcore Japan KK

GRAPHCORE

www.graphcore.ai



Neural network visualization from [POPLAR™](#)

GRAPHCORE

GRAPHCORE



2018 July 第一世代 IPU
Begins to ship the
PCI card.



IPU GC200



IPU M2000

2020 12月~ 第2世代IPU



IPU POD16 POD64/128/256



EXA POD



グラフコア・ジャパン株式会社
東京都千代田区
内幸町1-1-1
日比谷帝国ホテルタワー15F



CEO CTO
Nigel Toon Simon Knowles



2016 June Graphcore is officially founded as a company
2021 Jul over 620 employee, HQ Bristol, United Kingdom



Foundation
CAPITAL

pitango

CPVentures

Amadeus
Capital Partners

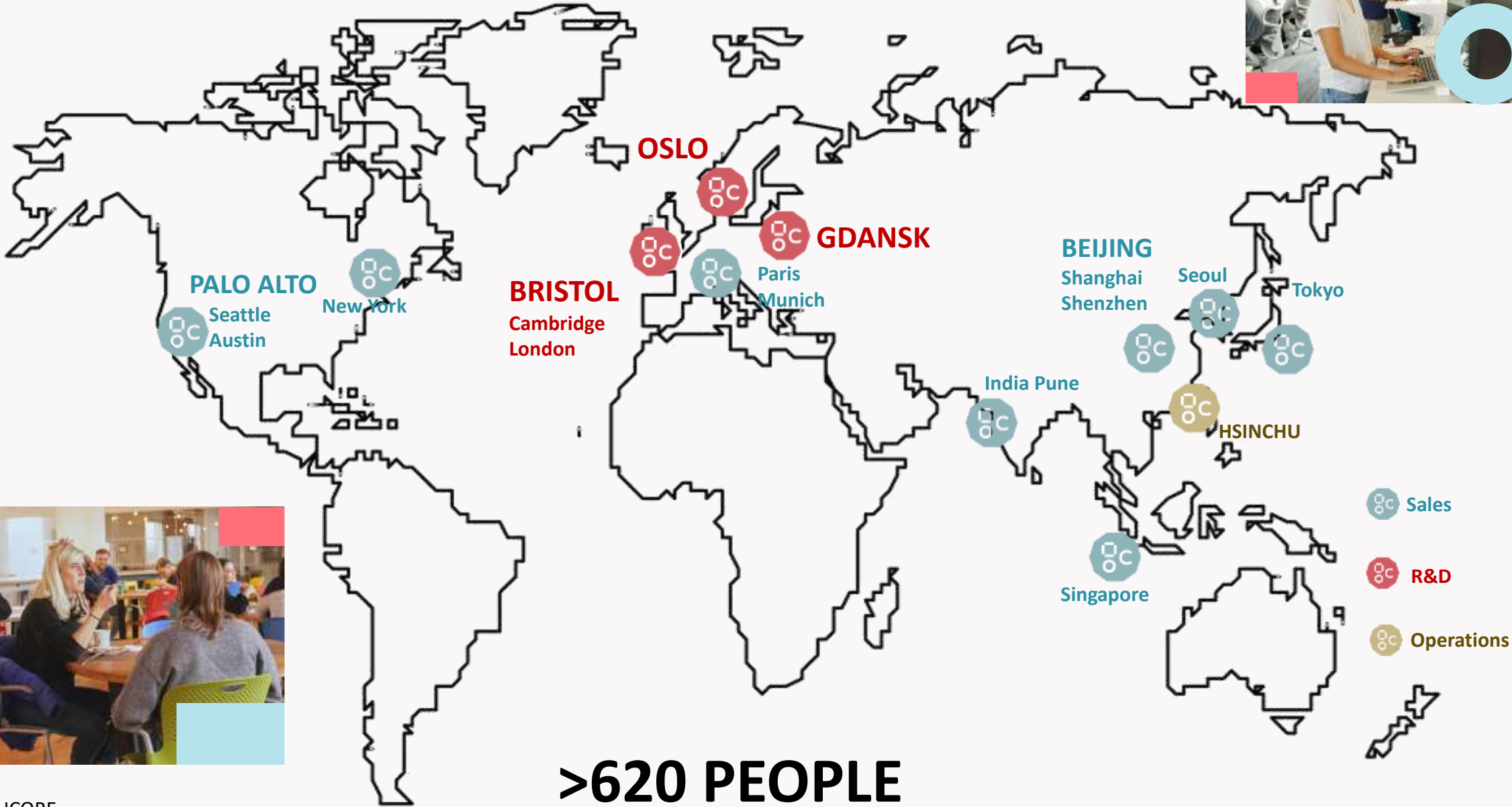
DELL
Technologies
CAPITAL

SAMSUNG
CATALYST
FUND

BOSCH
Robert Bosch Venture Capital GmbH

2016 Oct
Series A Funding
2020 Dec Series E
Total \$730M IN FUNDING

GRAPHCORE OFFICES



>620 PEOPLE

AGENDA

AI専用 革新プロセッサIPU

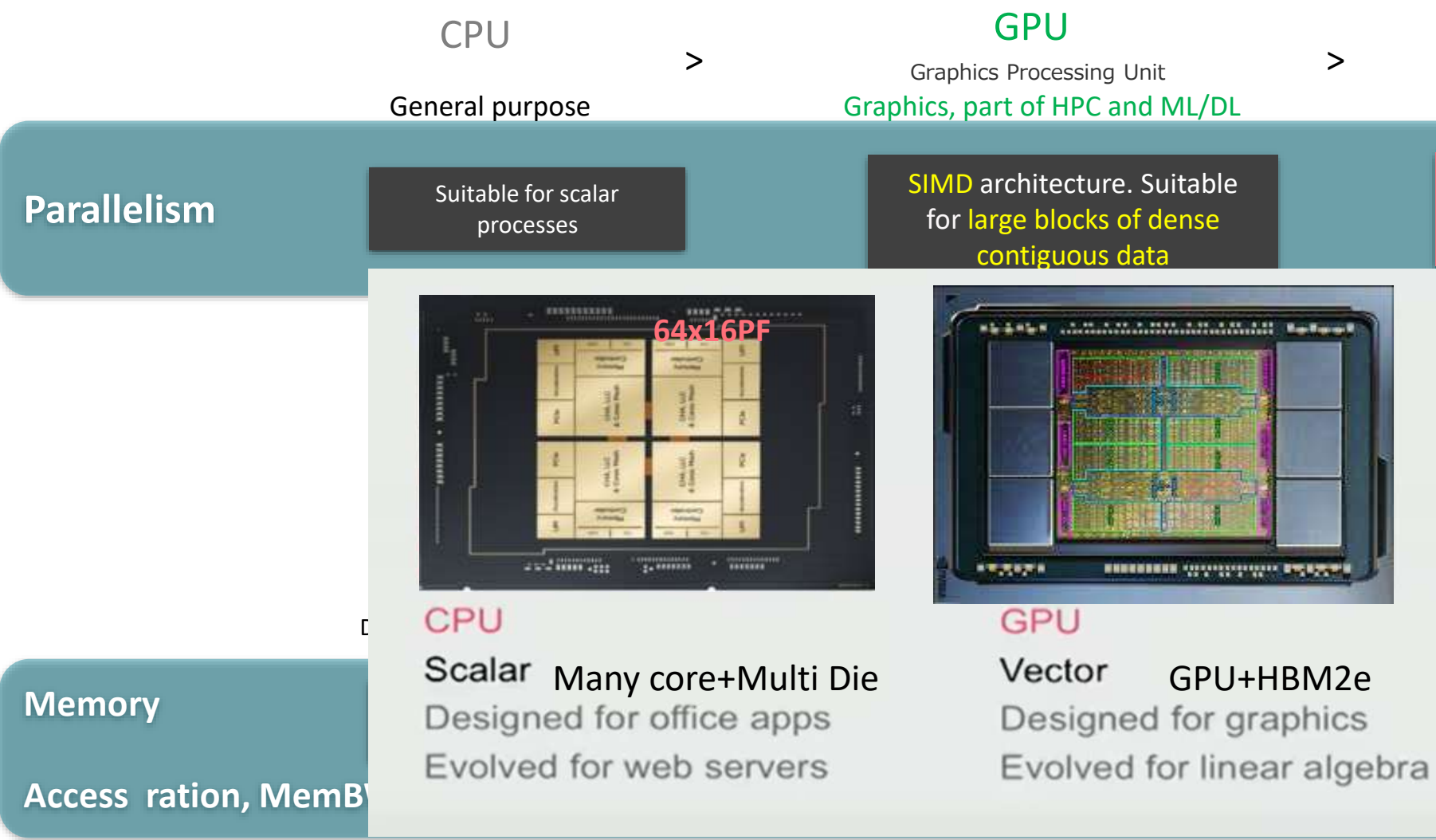
1. AI専用プロセッサ開発の背景
2. IPUおよびIPU-POD概要
3. 開発環境、PODPLAR SDKその他
4. USECASES, LARGE MODEL、EVAL/POC

<プログラム>

- 13:00-13:05 挨拶
佐藤三久（PCCC会長／HPC-OSS部会長／理研）
- 13:05-13:20 開催趣旨説明
滝沢寛之（東北大学）
- 13:20-13:45 「GPUとDPUによるアクセラレーテッドコンピューティング」
佐々木 邦暢(エヌビディア合同会社)
- 13:45-14:10 「AMD Instinct GPU and ROCm technical overview」
Greg, Oakes (AMD Sr. Solution Architect)
- 14:10-14:35 「AI革新プロセッサIPU」
中野 守（Graphcore Japan KK）
- 14:35-14:45 (休憩)
- 14:45-15:10 「インテルXPU戦略」
大内山 浩、高藤良史(インテル株式会社)
- 15:10-15:35 「ハイブリッドアクセラレータ SX-Aurora TSUBASA」
政岡靖久（日本電気株式会社）
- 15:35-16:00 「SambaNova Systemsのご紹介」
鯨岡俊則（SambaNova）
- 16:00-16:10 (休憩)
- 16:10-17:30 パネルディスカッション
「アクセラレータの王者と挑戦者たち」（仮題）
モデレータ：滝沢寛之（東北大学）



MACHINE LEARNING REQUIRES COMPUTE INNOVATION



Others

FPGA
TPU
IPU
:

AI REQUIRES COMPUTE INNOVATION BY IPU

A NEW APPROACH IS REQUIRED

AI ドメインスペシフィックな革新的アーキテクチャ

- AI本格的な応用段階、ML/DLの精度・機能・性能向上ニーズが急増
 - 大規模モデル パラメータ数は10年間で30万倍以上増加
 - 革新的アルゴリズムの台頭
 - TCO, 導入コスト, 電力コストを大幅に低減
- 今後10年間で1000倍以上の性能向上が必要

Parameters

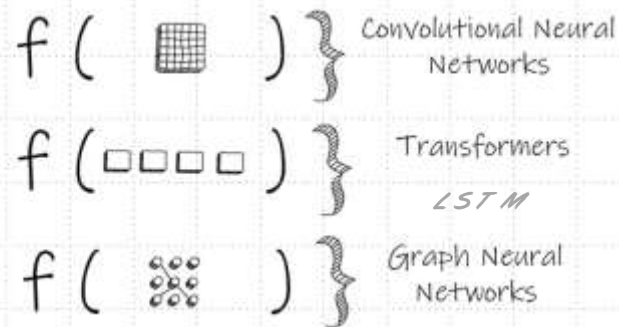
Massive parallelism

Sparsity in data structures

Low precision compute

Model parameter re-use

Static graph structure



● ResNet50: 25M
Jan 2016

● BERT-Large: 330M

Oct 2018

And more ...

GPT3: 175Bn

June 2020

OpenAI が開発

GPT-3 Generative Pretrained Transformer
1750億個のパラメータを使用した文章生成言語モデル. 通常の文章以外にも
ソースコードやデザインですらGPT-3」は理解し、生成する事が可能といわれている。

畳み込みニューラルネットワーク (CNN)

2016

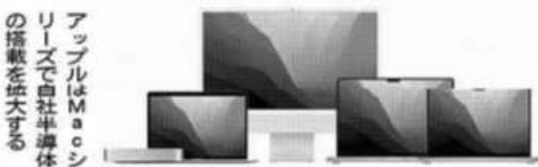
2019

Time

米IT、半導体開発に本腰

自前で半導体を開発する動きが広がる

IT大手	アップル(米)
	iPhoneやMacBookPro向けに自前で半導体を開発。5ナノメートルの技術を採用
自動車	グーグル(米)
	28日に発売するスマホ「ピクセル6」に独自半導体を採用。機械学習に特化した独自プロセッサも開発
通信	テスラ(米)
	19年から車載半導体の自社開発品に切り替え
	トヨタ自動車(日)
	集積チップやパワー半導体の開発を目的に20年にデンソーと新会社を設立
	シスコシステムズ(米)
	データセンター向けの機器などで自社開発の半導体を搭載
	ノキア(フィンランド)
	通信事業者向けのルーターに独自のプロセッサを搭載



アップルはMacシリーズで自社半導体の搭載を拡大する

世界のIT(情報技術)大手による半導体の開発競争が本格化している。米アップルは18日に発表したノートパソコンに自前で設計した半導体を採用する。グーグルも28日発売の新型スマートフォンに自社開発の半導体を搭載する。部品である半導体が製品そのものの競争力を左右するようになり、自動車や通信でも自前の開発能力を備える動きが広がる。

アップルは新型パソコン「MacBook Pro」向けに2種類の半導体で自社で設計開発した。2020年に発表した「M1」に続く独自半導体だ。処理性能は最大1.7倍で、動画にCG(コンピュータグラフィックス)を重ね合わせるような操作が容易になる。最先端の回路線幅5ナノ(ナノは10億分の1)の半導体製造プロセスで生産する。

グーグルも日本などで発売する「Pixel 6」に自社開発の半導体「Google Tensor」を採用し、画像処理や翻訳などの機能を高める。例えば、撮影した画像から不要な物体を消去できる「消しゴムマジック」と呼ぶ機能を搭載する。

自動車や通信もグーグルはこれまで米アルコムの半導体を使っていた。ハードウェアを担当するリック・オスという性格が強い。夏、日本経済新聞などの取材で「人工知能(AI)で新たな技術革新が起きている。汎用的な市販品では不十分と感じるよう

自社設計が競争力左右

アップル・グーグル、新商品に搭載

の頭脳である半導体に求められる性能は急速に高まっている。自社の用途に特化した専用半導体を作れば競合と差別化もできる。新たな使用環境や顧客ニーズに合わせた半導体をどれだけ早く開発し、製品に搭載できるかが競争能力を左右するようになりつつある。

IT大手が半導体に参入できるのは、開発と製造を水平分業するサプライチェーン(供給網)の存在が大きい。00年代に入り設計開発に特化したファブレスメーカーが台頭する一方、巨額の設備投資が必要な製造分野では台湾積体電路製造(TSMC)などの生産委託会社(ファブドリー)が製造技術を磨いてきた。さらにファウンドリーは半導体の設計機能も内製するようになった。半導体に知見の少ない企業も、独自に先端半導体を開発、搭載できるようになった。こうした自社設計の動きはアップルやグーグルにとどまらない。

自動車や通信も

グーグルはこれまで米アルコムの半導体を使ってきた。ハードウェアを担当するリック・オスという性格が強い。夏、日本経済新聞などの取材で「人工知能(AI)で新たな技術革新が起きている。汎用的な市販品では不十分と感じるよう

理由を説明している。これまでの半導体技術者を巡る獲得競争も過熱しており、企業規模で格差が開く可能性は否定できない。

東京大学の黒田忠広教授は「データや演算量の増加をつけ、半導体には高いエネルギー効率を求められている。十分な性能を引き出すには用途に合わせた設計が欠かせな

い」と指摘する。

コストの重い設計工程を効率化するため、東大はパナソニックやミライズテクノロジーズなどと共同プロジェクトを立ち上げた。省エネ半導体の開発効率を従来の10倍に高めるのが目標だ。

日本の半導体産業は1990年代から衰退が続

GRAPHCORE IPU (Intelligence Processor Unit)

AI Domain
Specific
Processor



IPU
GC200
(Processor)

TSMC 7nm 823mm² 59.4B Transistors
Many-cores 1472 tiles (Core, In-Proc-Mem) 8832 threads
In-proc Mem Total 900MB/IPU 47.5TB/s Mem BW
Peak Perf (fp16) 250TFlops AI-Float
(fp32) 62TFLOPS (fp32)

M2000

第2世代プロセッサ GC200を4基
1Uの筐体に搭載したAIアクセラレータ



筐体をIPU-Linkで相互接続することで、性能をスケールアップ
IPU環境を容易に使うための POPLAR SDK を提供
PyTorch や TensorFlow 等の主要なフレームワークをサポート



IPU-POD256

IPU-POD Family



IPU-POD16

IPU-POD64

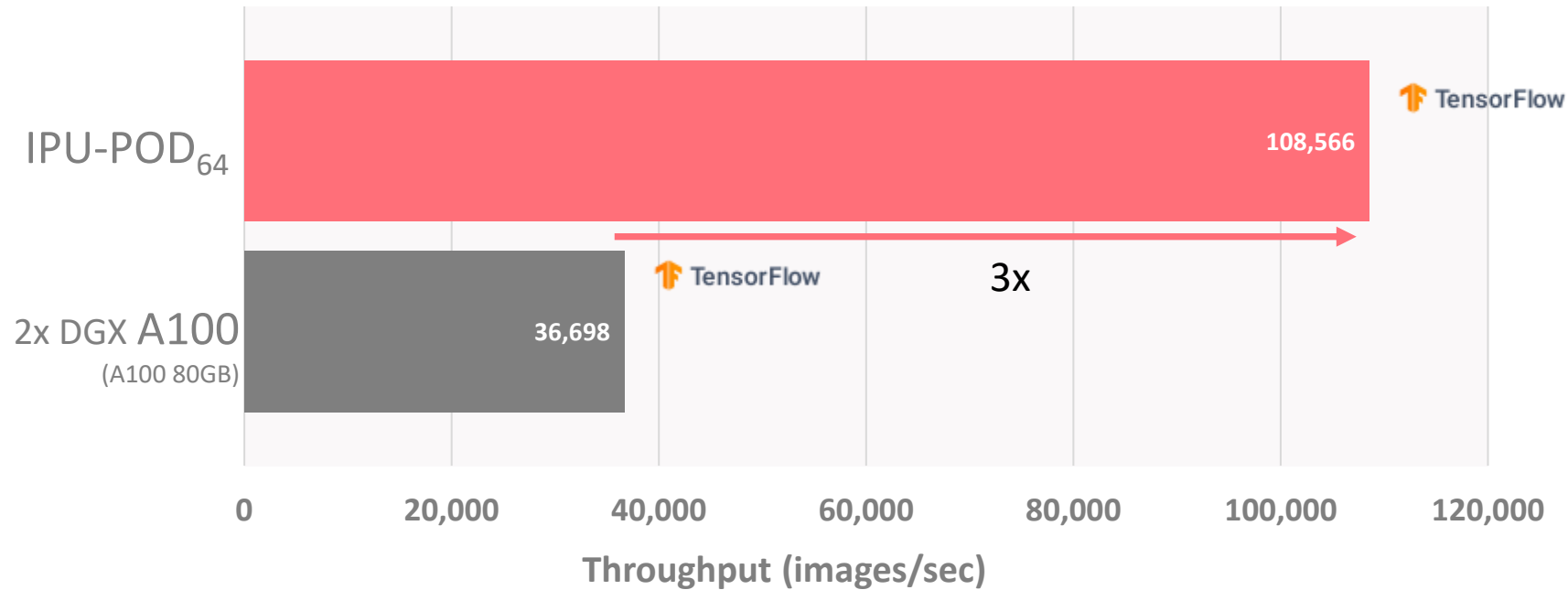
- 1筐体で 1ペタFlops AI-Float(fp16)



CV RESNET50 TRAINING

3x Higher Throughput | 3.5x Perf/\$ Advantage

同レベルのシステム価格の二つのシステム、IPUとGPUの性能を比較



NOTES: (updated 20th Dec 2021)

ResNet-50 v1.5 Training Highest Throughput | ImageNet2012 Dataset

IPU-POD64 (16x IPU-M2000) | FP 16.16 | SDK 2.4.0 | TensorFlow | <https://www.graphcore.ai/performance-results>

Best results for DGX A100 640GB (A100-SXM4-80GB) using TensorFlow | Mixed Precision | using 0.90x scaling from 1=>2 DGX A100

DGX A100 results published by NVIDIA (<https://developer.nvidia.com/deep-learning-performance-training-inference>) on 1 Dec 2021

SUMMARY

MODEL NAME:

RESNET-50 V1.5

ML DOMAIN:

COMPUTER VISION

MODEL DESCRIPTION:

RESNET IS A LEGACY IMAGE CLASSIFICATION MODEL, OFTEN USED AS A PROXY FOR INITIAL PERFORMANCE EVALUATION

DEPLOYMENT:

TRAINING

FRAMEWORK:

TENSORFLOW 1
[TF2 ALSO IN MODEL GARDEN]
[PYTORCH ALSO IN MODEL GARDEN]

IPU ADVANTAGE:

IPU-POD64 ACHIEVES 3X HIGHER THROUGHPUT VS TWO DGX A100, DELIVERING >3.5X PERF/\$ ADVANTAGE



Pricing Assumptions

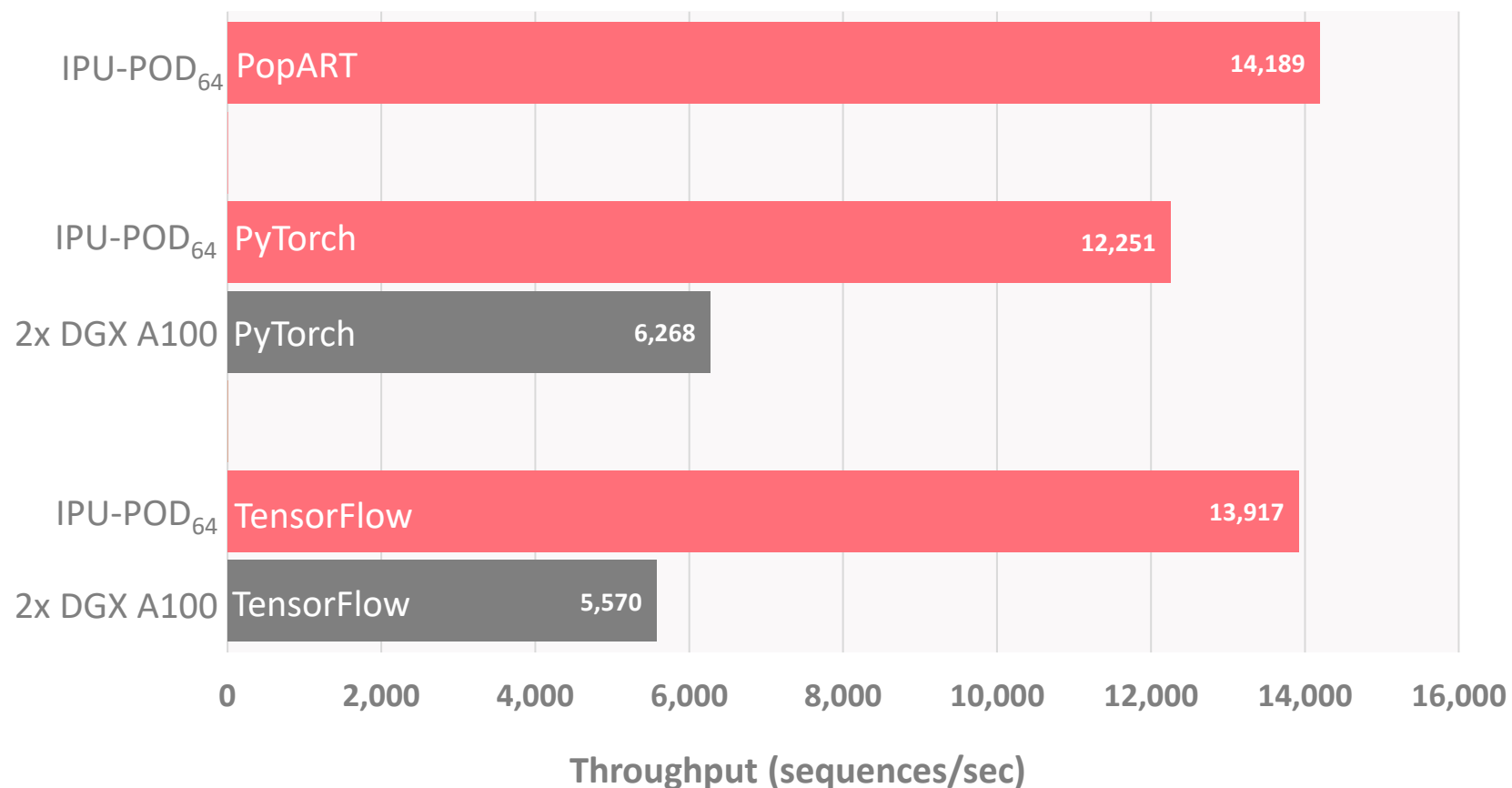
MSRP

IPU-POD64 (4-host)	~\$500k
2x DGX A100 640GB	~\$600k

NLP BERT-LARGE PRE-TRAINING

2.3x Higher Throughput | 2x Perf/\$ Advantage

同レベルのシステム価格の二つのシステム、IPUとGPUの性能を比較



NOTES: (updated 20th Dec 2021) | headline speedup based on TensorFlow vs TensorFlow throughput comparison
BERT-Large Phase 1 Pre-Training Throughput (SL128)
IPU-POD64 (16x IPU-M2000) | SDK 2.4.0 | <https://www.graphcore.ai/performance-results>
DGX A100 320GB (A100-SXM4-40GB) | Mixed Precision | Using 97% perf scaling based on published BERT results
DGX A100 PyTorch <https://developer.nvidia.com/deep-learning-performance-training-inference> 1st Dec 2021
DGX A100 TensorFlow - https://ngc.nvidia.com/catalog/resources/nvidia:bert_for_tensorflow/performance

SUMMARY

MODEL NAME: BERT-LARGE

ML DOMAIN: NLP

MODEL DESCRIPTION:
BERT IS AN UNSUPERVISED TRANSFORMER-BASED LANGUAGE REPRESENTATION MODEL. IT IS ONE OF THE MOST WIDELY USED NLP MODELS

DEPLOYMENT:
PRE-TRAINING PHASE 1 (SL128)

FRAMEWORK:
TENSORFLOW 1
PYTORCH
POPART

IPU ADVANTAGE:
IPU-POD64 ~2.3X HIGHER THROUGHPUT THAN 2X DGX A100 THROUGHPUT, DELIVERING ~2X PERF/\$ ADVANTAGE

Pricing Assumptions

MSRP

IPU-POD64 (1-host)	~\$440k
2x DGX A100 320GB	~400k

GRAPHCORE



GRAPHCORE MLPERF V1.1 RESULTS FOR BERT LARGE



MLPERF 1.1 RESULTS

DECEMBER 2021

既存モデルをIPUアーキテクチャに最適化を施した事例

IPU: トランスフォーマーモデルの最適化

2x – 4x 学習時間の高速化 | 3x モデルサイズの縮小

Graphcore Research Grの研究論文

Group Transformer

IPUアーキテクチャを活かしグループ化された
畳み込みモジュールを追加

3x パラメータ数の縮小

2x 学習時間の高速化

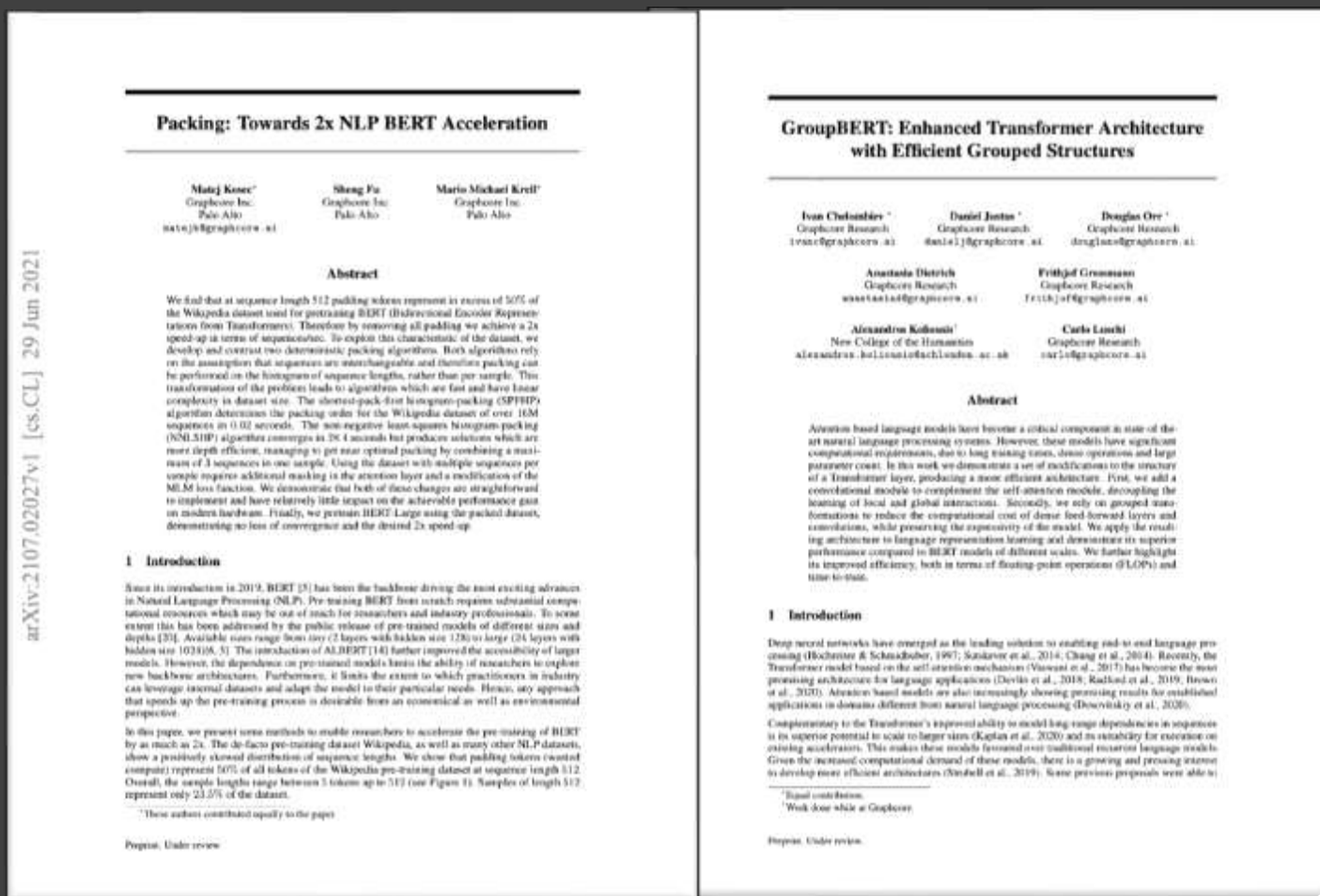
結果の精度向上

Transformer Packing

データパディングの除去（IPUではなくGPUで
必要） 非負の最小二乗法によるヒストグラムの
パッキングが 計算効率を大幅に向上。学習時間を
2倍に短縮

<https://arxiv.org/abs/2107.02027>

<https://arxiv.org/abs/2106.05822>

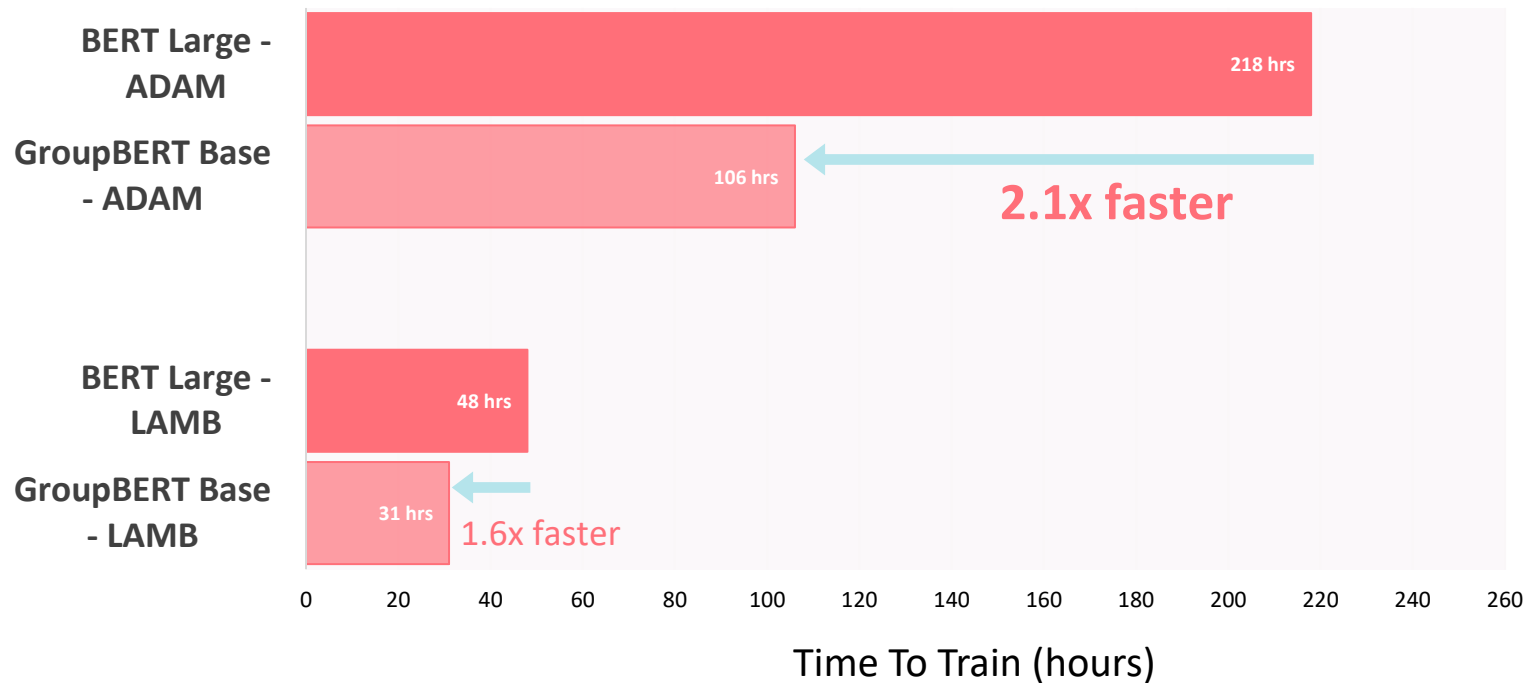


GRAPHCORE

NLP GroupBERT

leverages IPU architecture capabilities
adds grouped convolution module
2x reduction in # parameters
2x faster time to train
improved accuracy of results

GroupBERT Base vs BERT Large Time-to-accuracy



NOTES: (updated 20th Dec 2021)

SUMMARY

MODEL NAME:

GROUPBERT

ML DOMAIN:

NLP

MODEL DESCRIPTION:

GROUPBER IS AN OPTIMISED
IMPLEMENTATION OF THE BERT
MODEL THROUGH USE OF
GROUPED TRANSFORMATIONS

DEPLOYMENT:

TRAINING

FRAMEWORK:

TBC

IPU ADVANTAGE:

GROUPBERT DELIVERS 2X REDUCTION IN
NUMBER OF PARAMETERS AND 2X FASTER
TIMETO TRAIN FOR EQUIVALENT MODEL
PERFORMANCE VS STANDARD BERT



<https://arxiv.org/abs/2106.05822>

MACHINE LEARNING REQUIRES COMPUTE INNOVATION



CPU

GPU

IPU

Intelligence Processing Unit
ML/DL

General purpose

Graphics Processing Unit
Graphics, part of HPC and ML/DL

Parallelism

Suitable for scalar processes

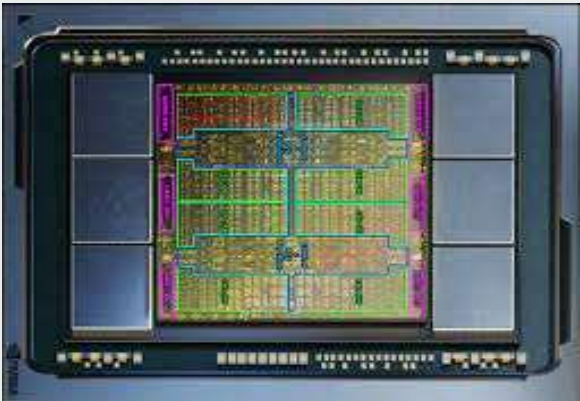
SIMD architecture. Suitable for large blocks of dense contiguous data

Massively parallel MIMD. High performance/efficiency as ML trends to sparsity & small kernels



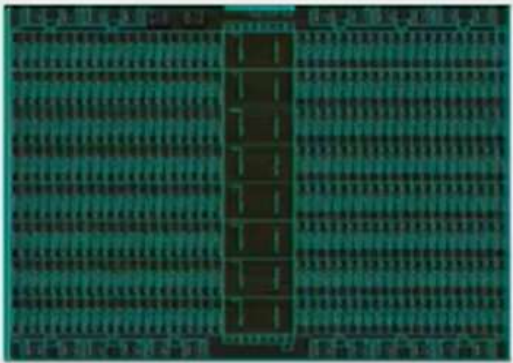
CPU

Scalar Many core+Multi Die
Designed for office apps
Evolved for web servers



GPU

Vector GPU+HBM2e
Designed for graphics
Evolved for linear algebra



IPU

Graph
Designed for intelligence

Memory

Access ration, MemB

Processor

Memory

GC200 IPU PROCESSOR

IPU-Tiles™

1472 independent IPU-Tiles™ each with an IPU-Core™ and In-Processor-Memory™

IPU-Core™

1472 independent IPU-Core™

8832 independent program threads executing in parallel

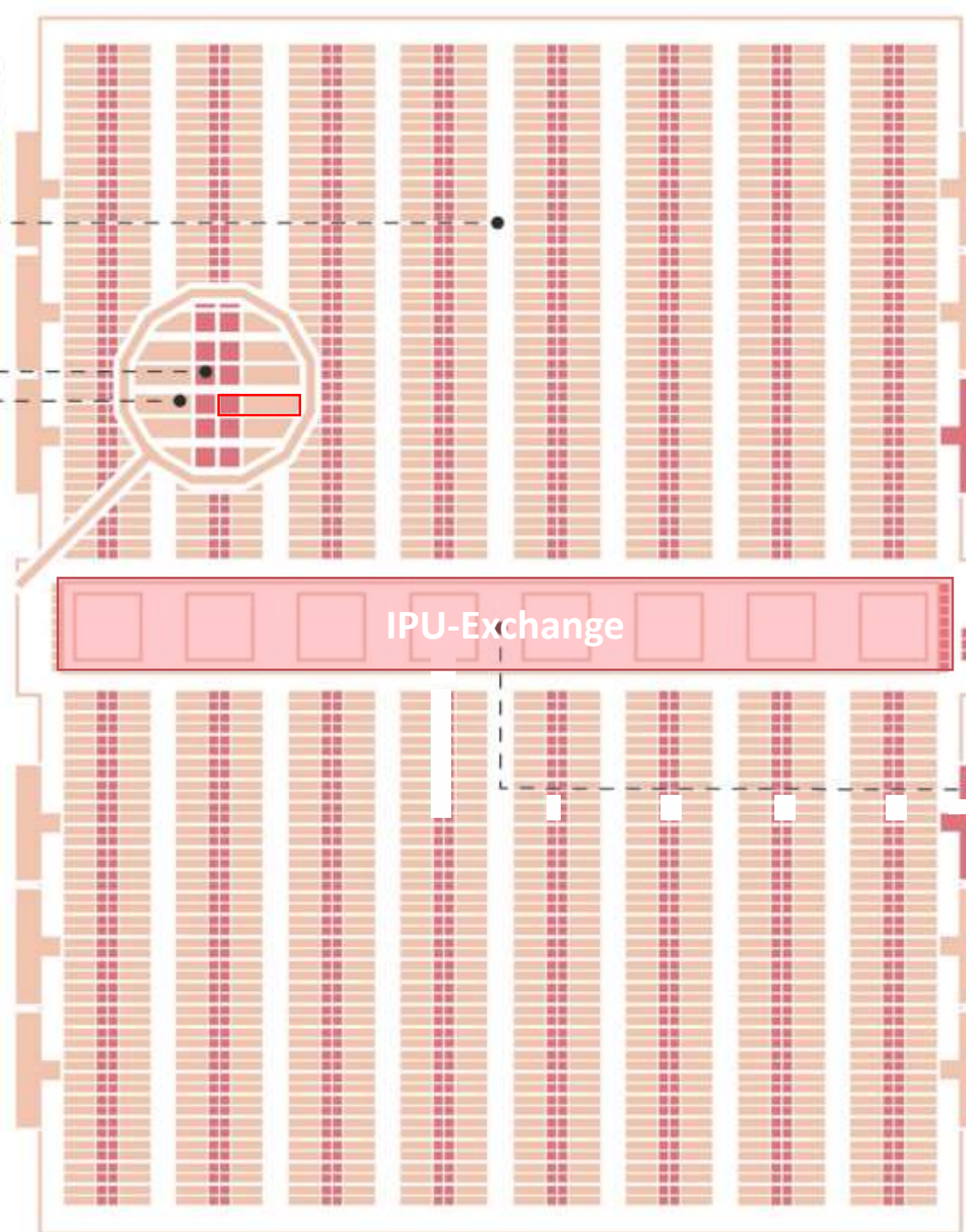
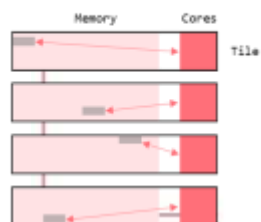
In-Processor-Memory™

900MB In-Processor-Memory™ per IPU

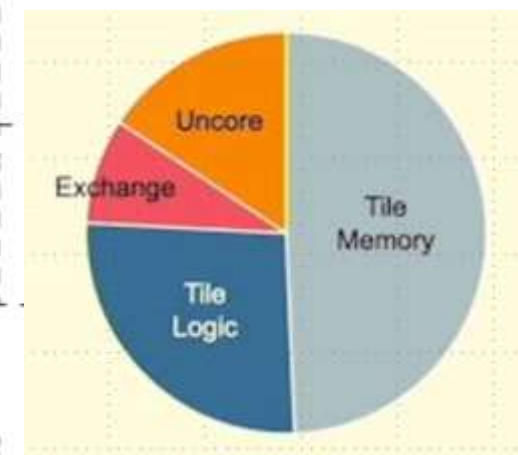
47.5TB/s memory bandwidth per IPU

- **IPU Tiles** 1472 (Core+Mem)
8832 threads
- **in-Processor Memory**
900MB per IPU (620KB/tile)
47.5TB/s Memory BW per IPU
(33GB/s per tile)
- **IPU-Exchange**
8 TB/s All to All
- **PCIeG4x16** x2

IPU's distributed SRAM



GC200 TSMC 7nm @ 823mm²
250TFlops AI-Float (fp16)
@ 1.325 GHz



10 x IPU-Links,
320GB/s chip to chip bandwidth

IPU-Machine: M2000

4 x Colossus™ GC200 IPU
1 petaFLOPS AI compute
Up to 526GB Exchange Memory™
2.8Tbps IPU-Fabric™

Each Colossus™ GC200 IPU

59.4Bn transistors, TSMC 7nm @ 823mm2
250 teraFLOPS AI compute
1472 independent processor cores
8832 separate parallel threads

IPU-Gateway SoC

Arm Cortex-A quad-core SoC
Super low latency IPU-Fabric™ interconnect

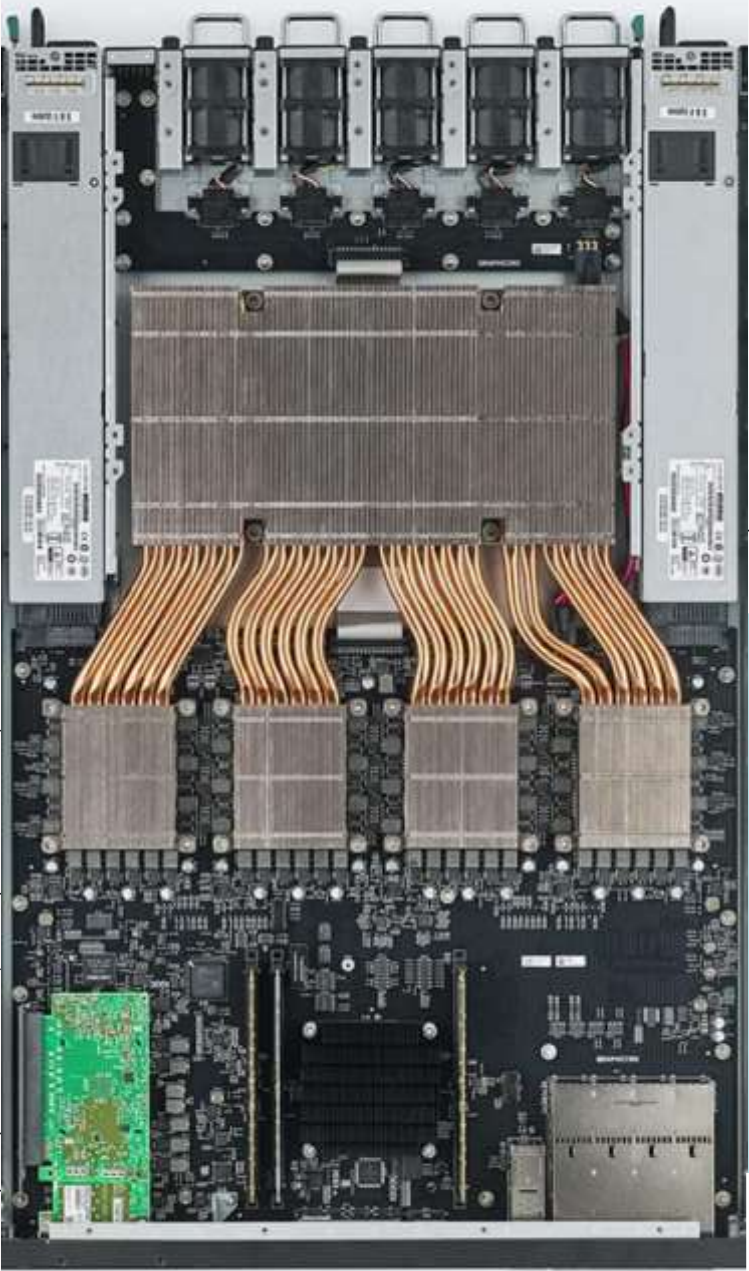
M.2 Connector

Board Management Controller

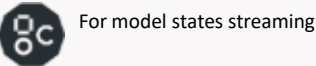
M.2 Slot

PCIe FH3/4L G4x8 Slot (RNIC/SmartNIC)

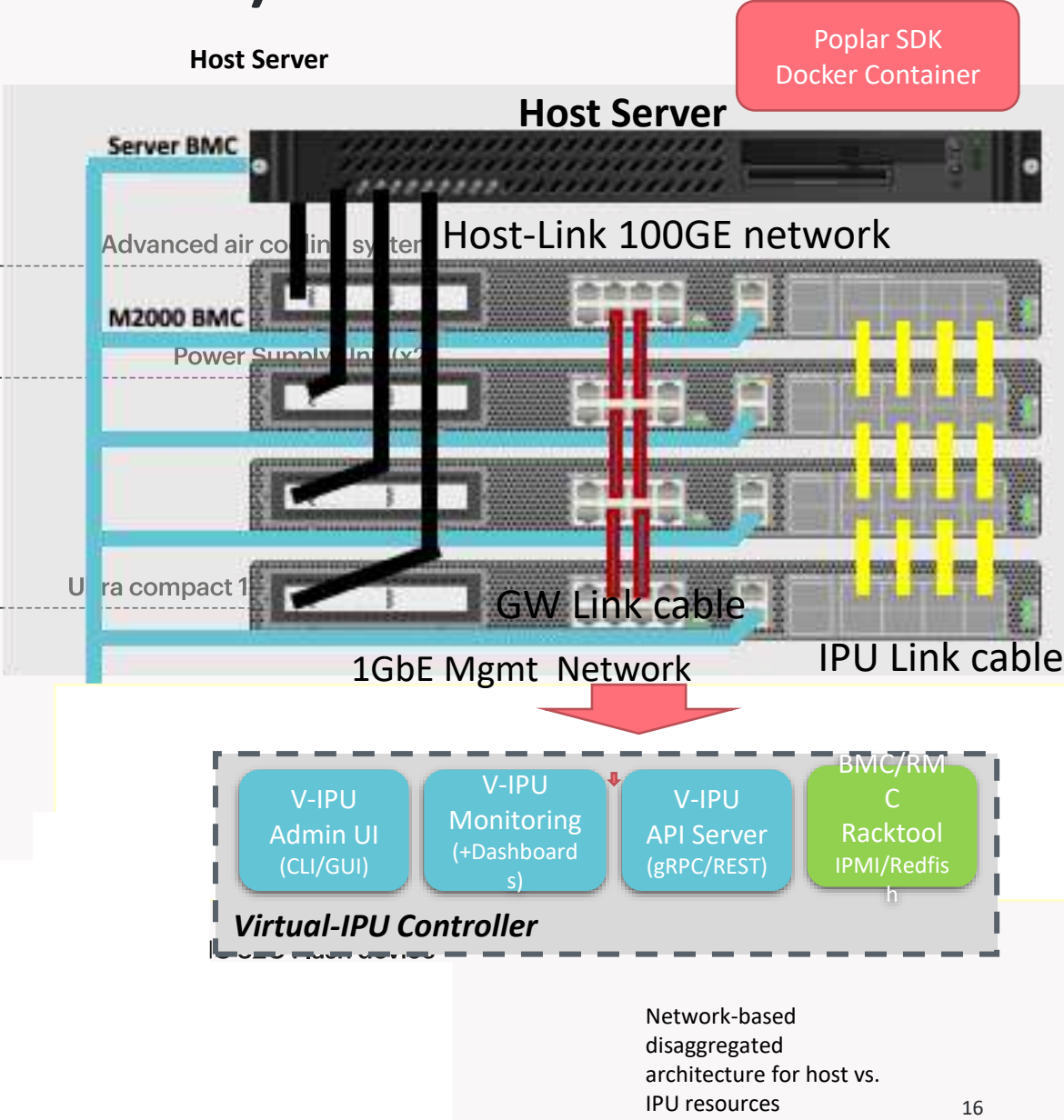
DDR4 DIMM DRAM x 2



Power
Power consumption 900W-1100W (typical)
Power cut(Max) (1500W)
Input Power 94-256Vac (115-230VAC Normal)



EX) POD16 DA SYSTEM



BSP COMPUTE SET EXECUTION

Bulk synchronous parallel

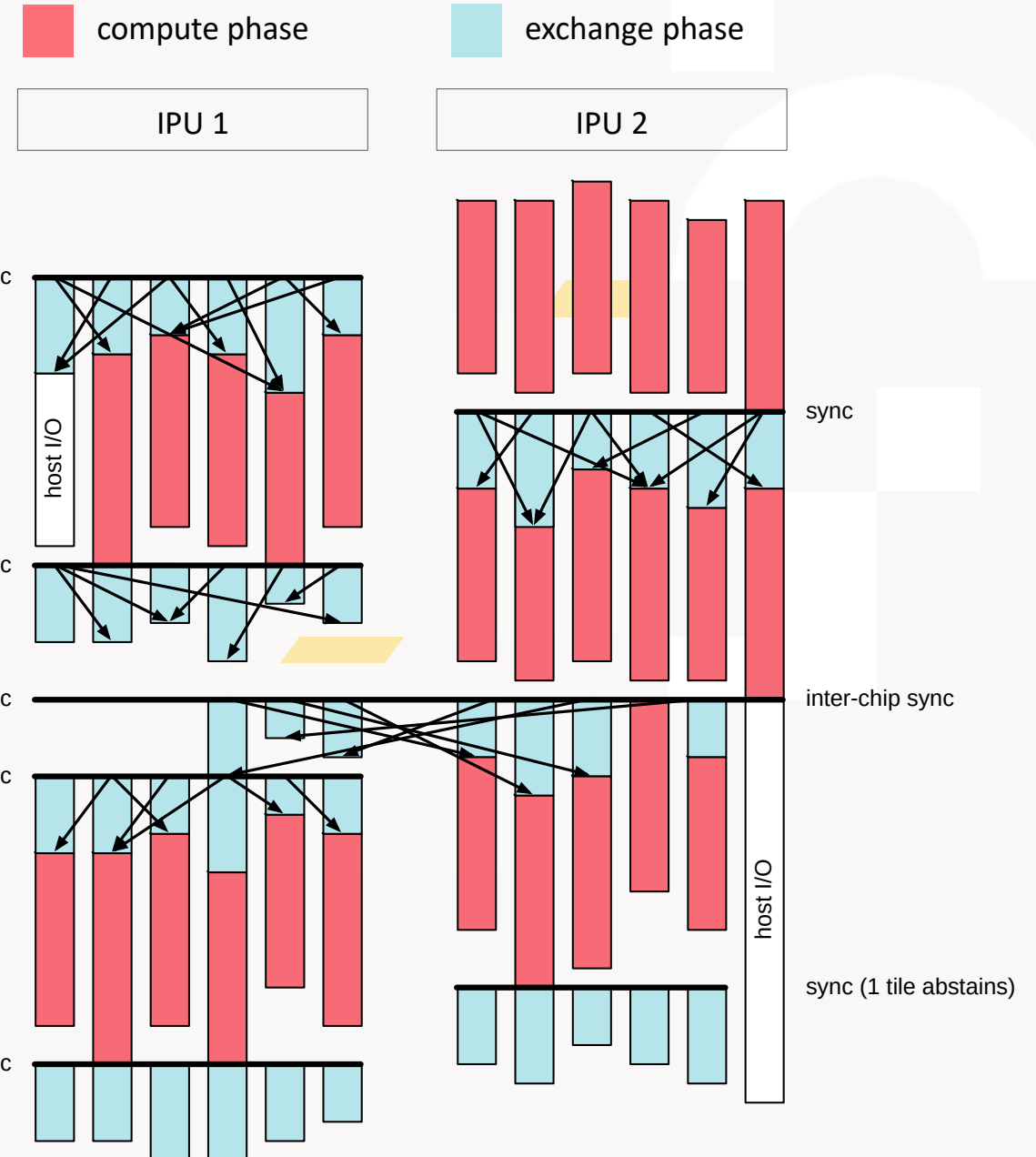
Bulk Synchronous Parallel (BSP)

compute | synchronize | exchange

BSP Compute Set execution works across multiple IPU devices

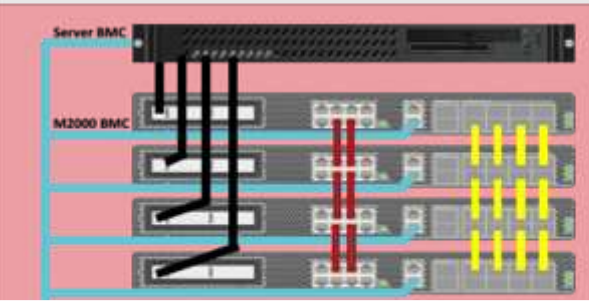
Inter- and Intra-IPU sync phases ensure no live/dead locks or race conditions

Tiles exchanging data with the host can abstain from syncing with other tiles



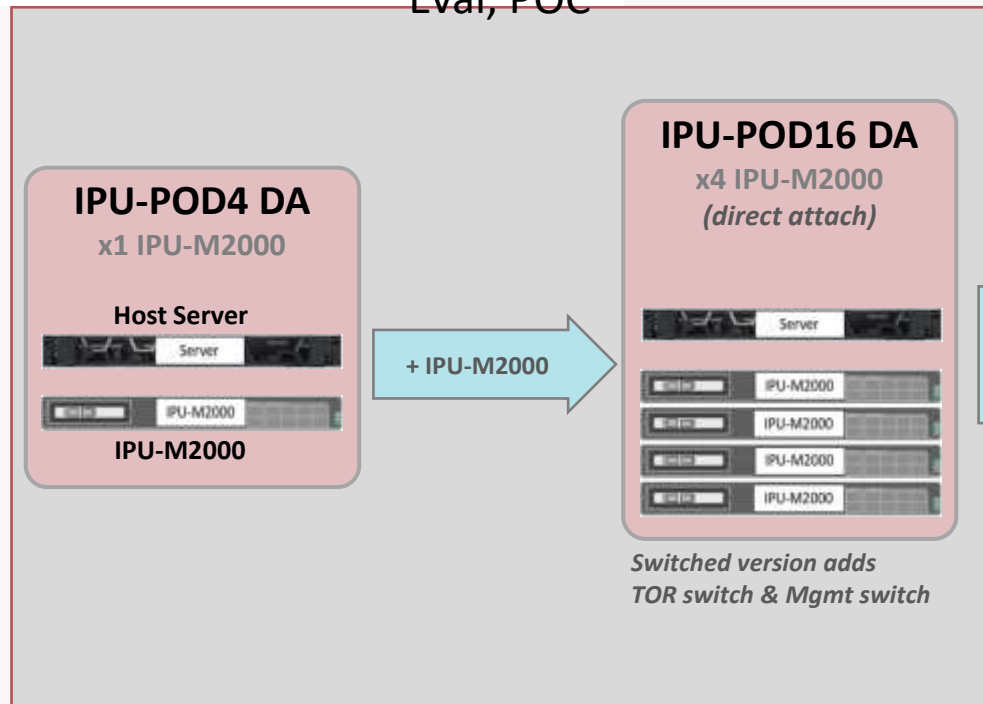
IPU-POD FLEXIBLE SCALE OUT SYSTEM

Reusable, simple migration between platforms



IPU-POD16 DA (Direct attached)

Eval, POC



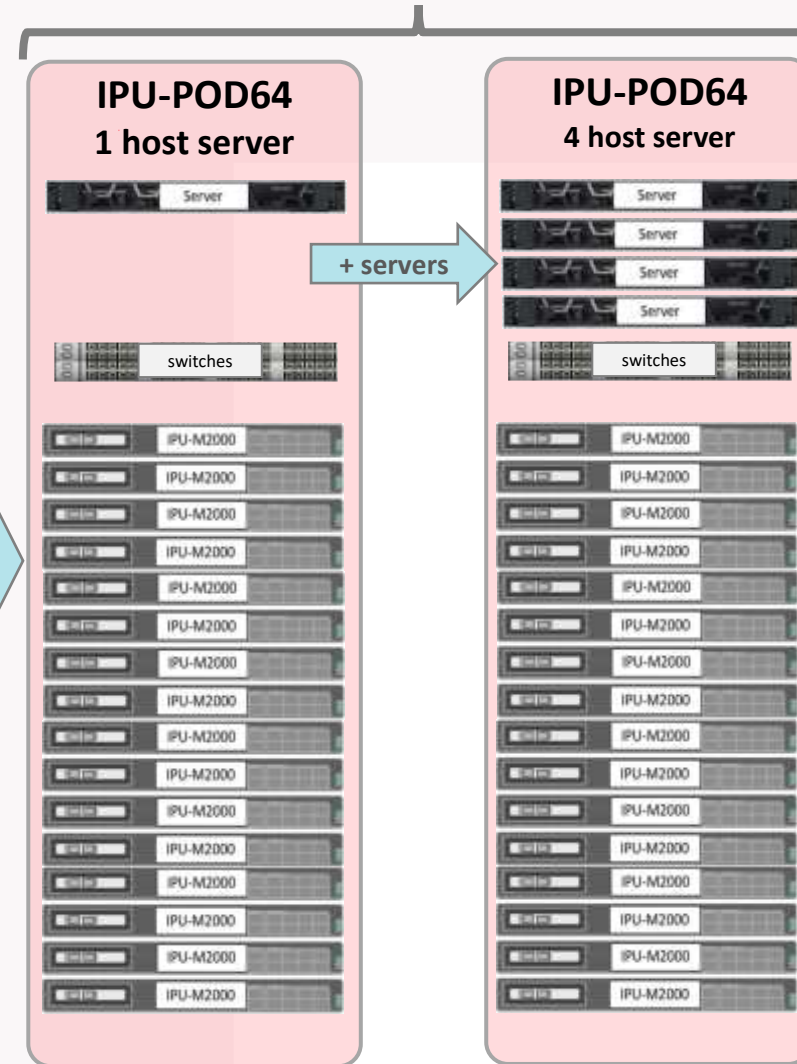
Direct Attached, Pre-configured Model
POD4,16,32



1x IPU-M2000

4x IPU-M2000

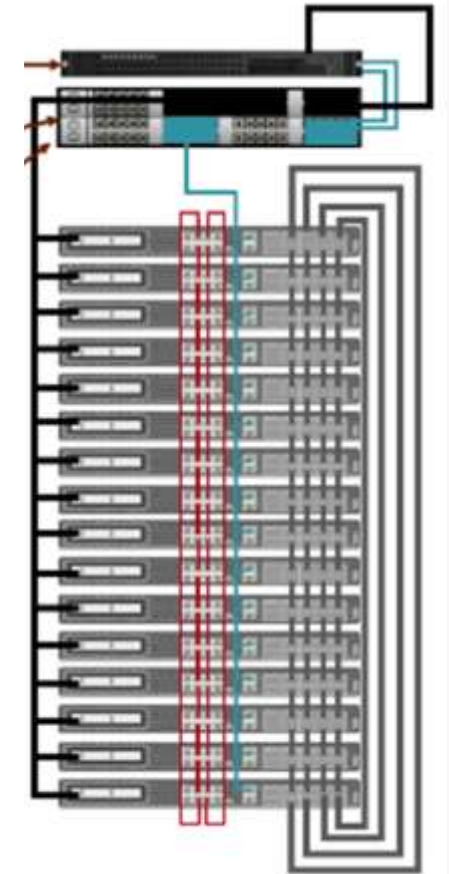
Disaggregated Server enables configurable Host Server /IPU ratio depending on workload for optimised TCO



16x IPU-M2000

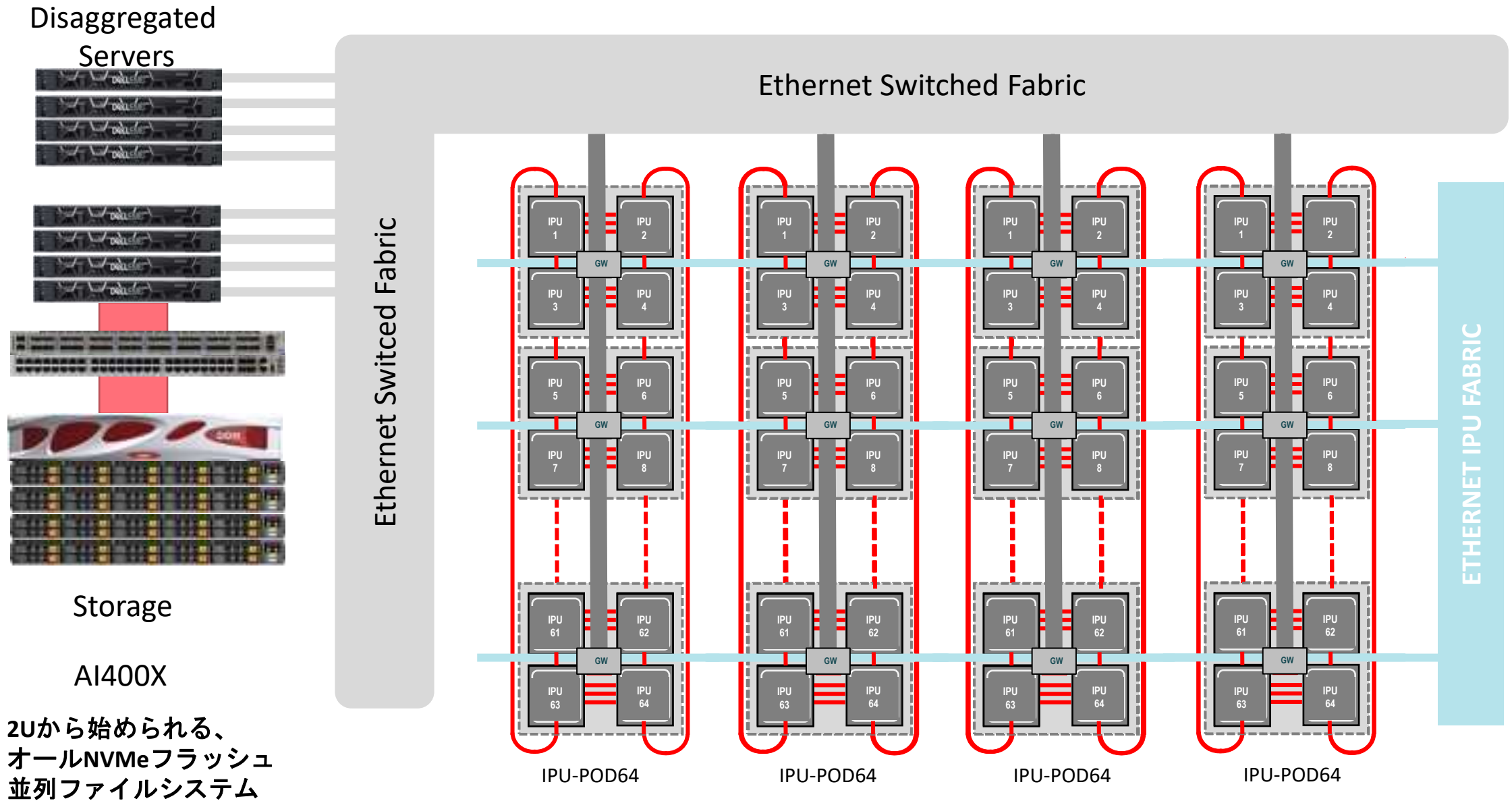
Production

IPU-POD64
(Scale out)



Upto 16Kx IPU-M2000

IPU-POD256 DATACENTER ARCHITECTURE



IPU POD SYSTEM



EXPLORE

IPU-POD₁₆

- 4x 1U IPU-M2000
- 1x dual-CPU server
- 4 PetaFLOPS AI compute



BUILD

IPU-POD₆₄

- 16x 1U IPU-M2000
- 1x dual-CPU server for BERT
- 4x dual-CPU server for ResNet
- 16 PetaFLOPS AI compute



GROW.1

IPU-POD₁₂₈

- 32x 1U IPU-M2000
- 2x dual-CPU server for BERT
- 8x dual-CPU server for ResNet
- 32 PetaFLOPS AI compute



GROW.2

IPU-POD₂₅₆

- 64x 1U IPU-M2000
- 4x dual-CPU server for BERT
- 16x dual-CPU server for ResNet
- 64 PetaFLOPS AI compute



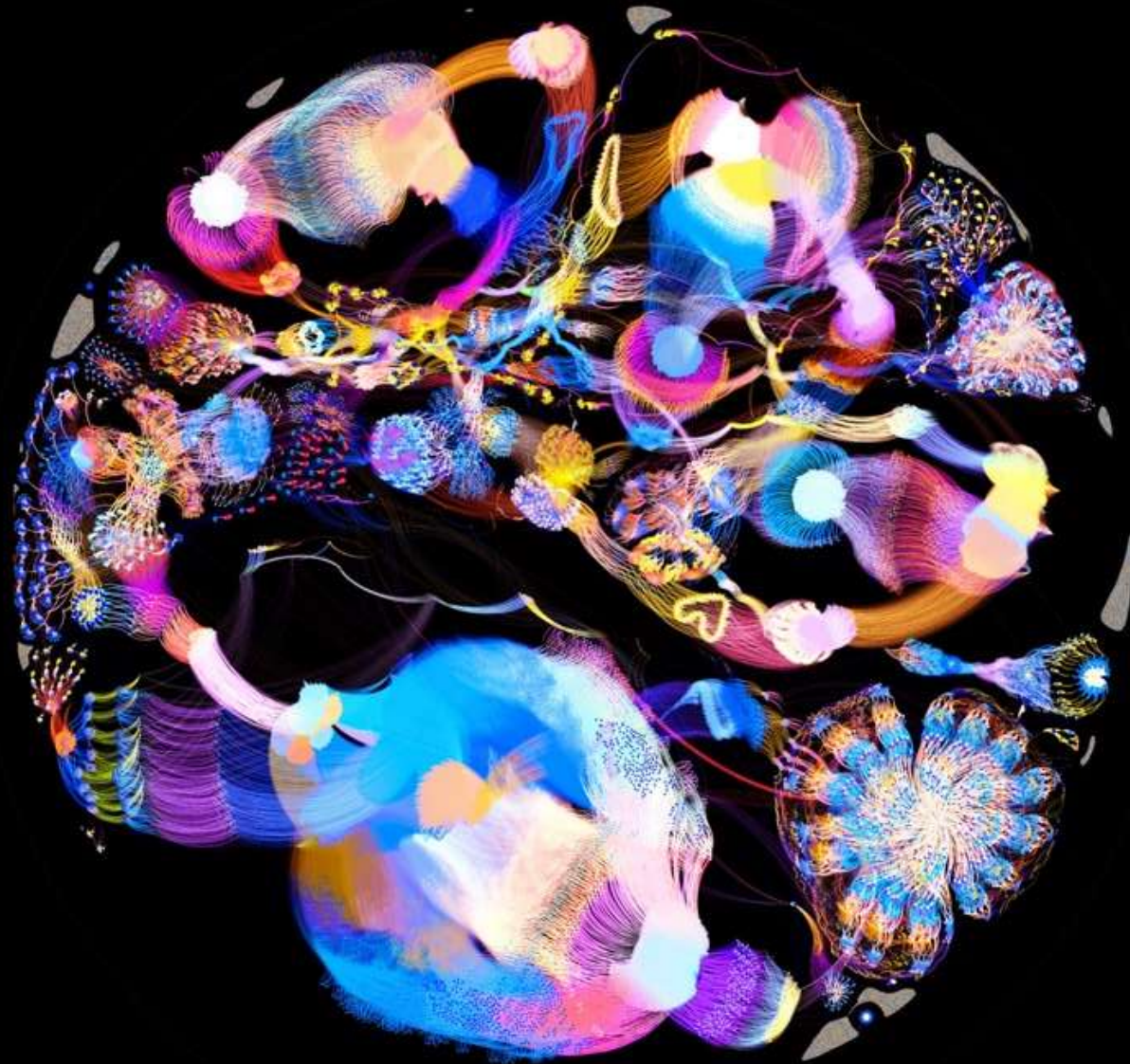
Large

IPU-POD_{512, 1024} EXA-POD



SOFTWARE

For Ease of Use



POPLAR™ COMPUTE GRAPH VISUALISATION

POPLAR SDK

Host Server OS



ubuntu

18.04

20.04 (preview)



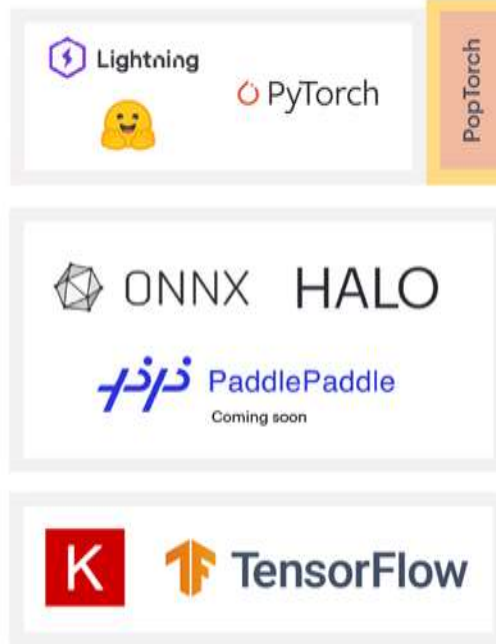
CentOS

7.6

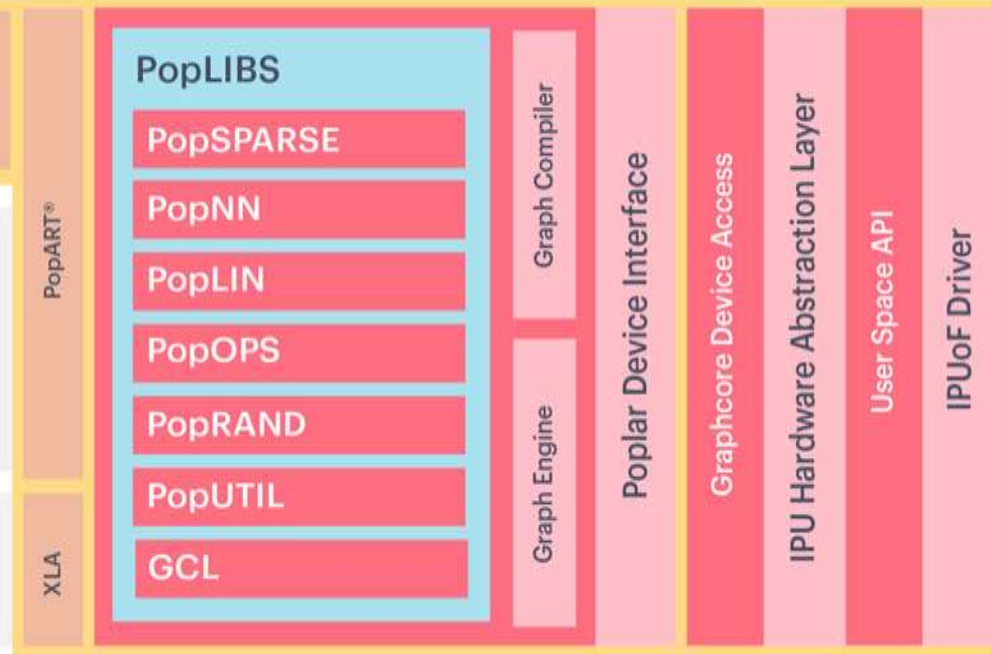
(RHELへの移行)

Debian 10

ML Frameworks



POPLAR®



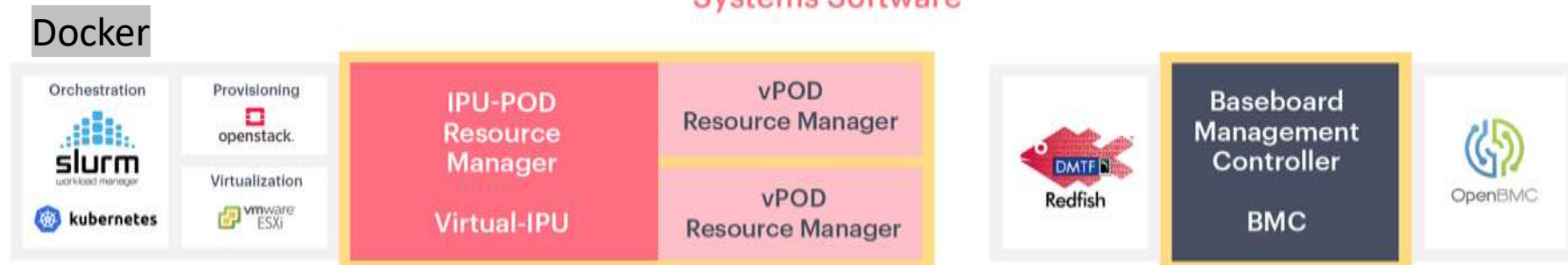
Drivers

Platform Hardware

IPU-POD Systems



Systems Software



POPLAR

Visualization
& Debug Tools



Management data can be delivered via Redfish DMTF to management systems such as Grafana.
Poplar software and our IPU support containerization with Docker.



CLOUD NATIVE & SCALEABLE

SOFTWARE DEFINED INFRASTRUCTURE

Development



POPLAR

Orchestration &
Scheduling



kubernetes



Monitor &
Control



Redfish



OpenBMC

Logging &
Visualization



Prometheus

Provisioning &
Virtualization

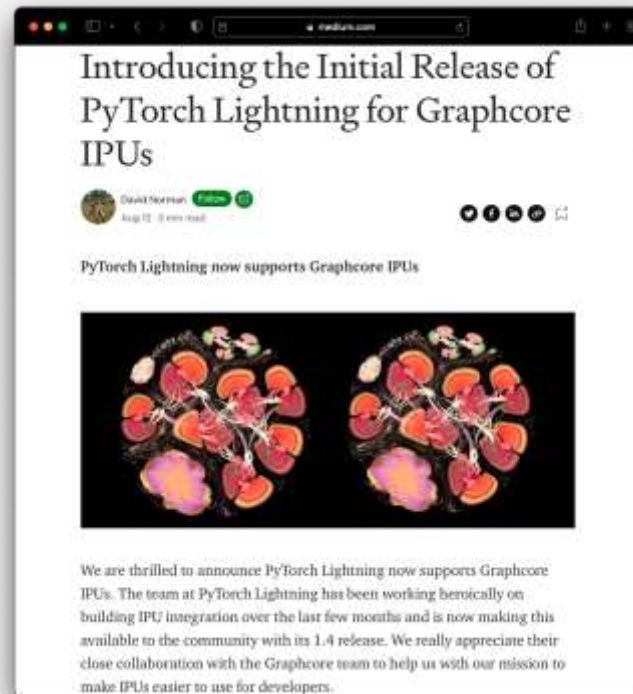
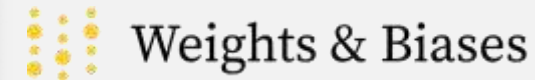
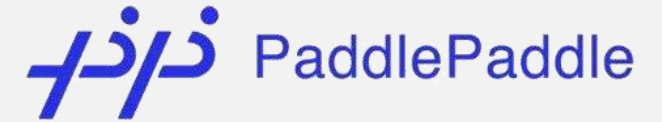
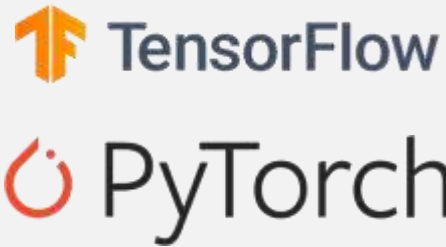


Virtual-IPU



GRAPHCORE

DEVELOPER ECOSYSTEM WITH PARTNERS



GRAPHCORE

GRAPHCORE CONFIDENTIAL

RESOURCES CENTRE

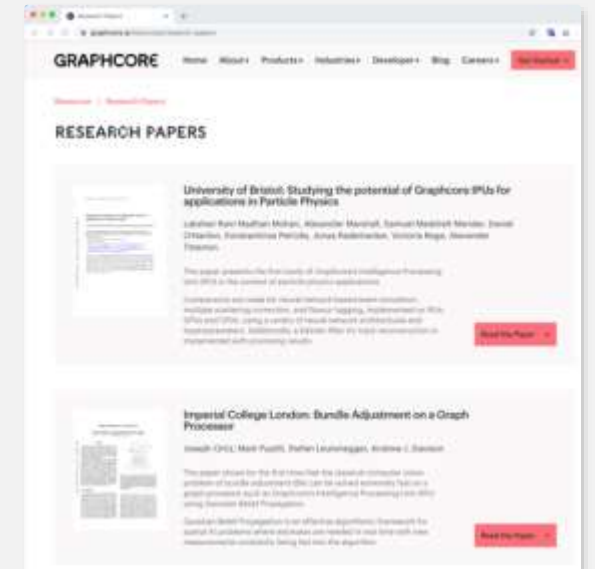
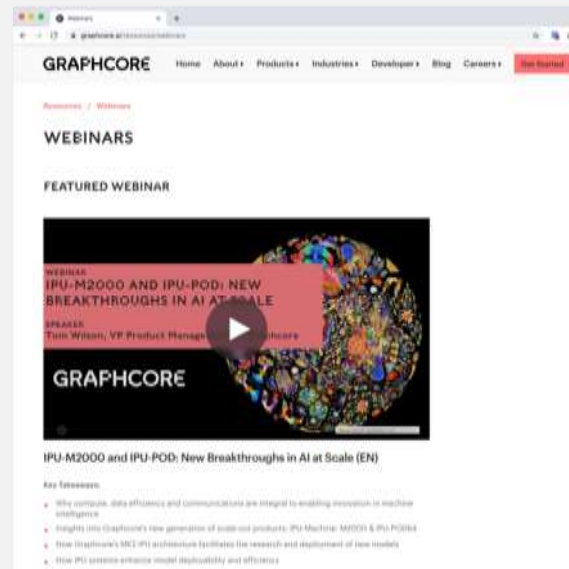
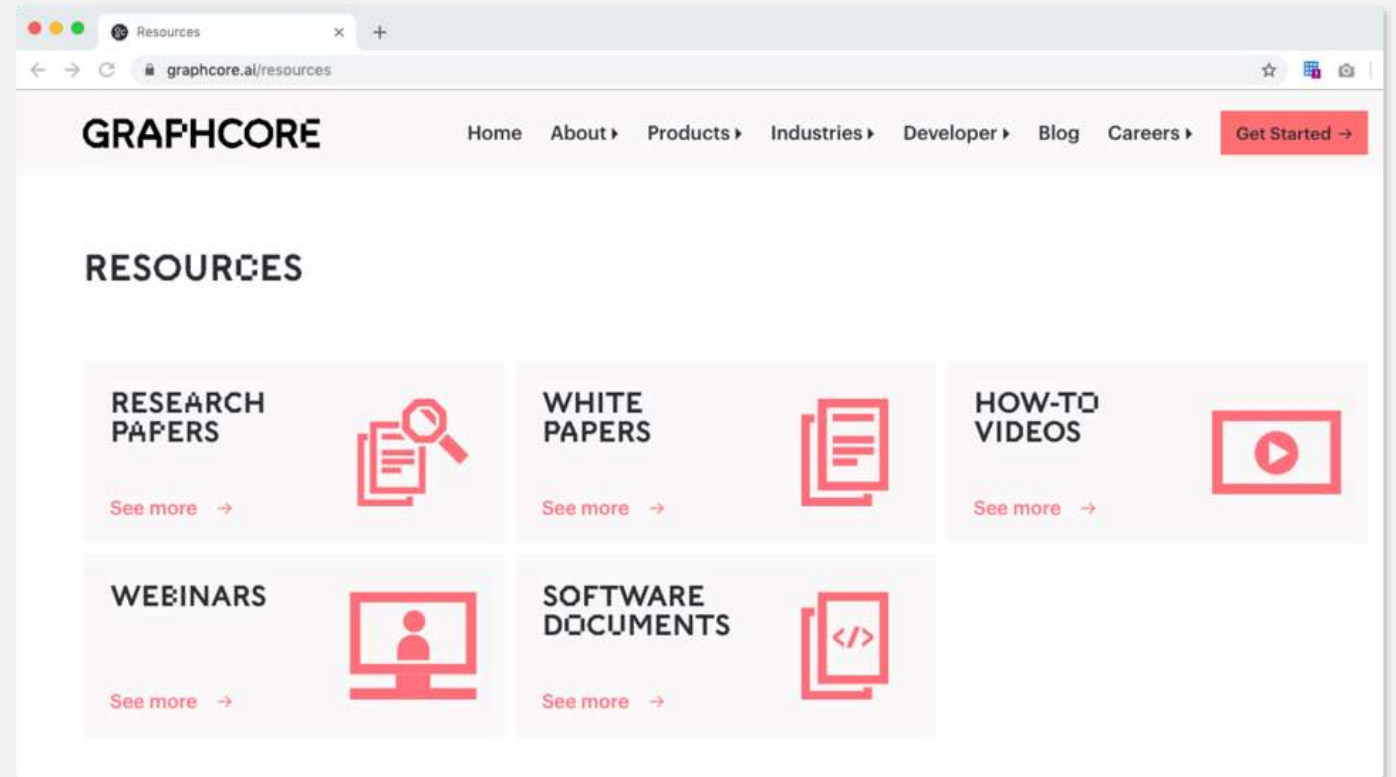
<https://www.graphcore.ai/resources>

- Central source of research papers, white papers, how-to videos, on-demand webinars and documentation

NEWSFLASH: GRAPHCORE OPEN SOURCES POPLAR GRAPH LIBRARIES (POPLIBS)

<https://www.graphcore.ai/posts/graphcore-open-sources-poplar-graph-libraries-poplibs>

GitHub



RAPIDLY INCREASING NUMBER OF SERVER PLATFORM PARTNERSHIPS...

The Dell logo, consisting of the word "DELL" in a blue, sans-serif font with a diagonal slash through the "E".The Hewlett Packard Enterprise logo, featuring a green rectangular icon above the text "Hewlett Packard Enterprise" in a black, sans-serif font.The Inspur logo, featuring the word "inspur" in a blue, italicized, sans-serif font.The Atos logo, featuring the word "Atos" in a blue, sans-serif font.The NEC logo, featuring the letters "NEC" in a bold, blue, sans-serif font.The Supermicro logo, featuring the word "SUPERMICR" in a blue, sans-serif font with a red dot, enclosed within a green oval.The 2crsi logo, featuring the word "2crsi" in a bold, sans-serif font with an orange "2" and "cr" in orange, and "si" in black.



STORAGE PARTNERS

**GRAPHCORE**
ELITE PARTNER 



GLOBAL ENGAGEMENT WITH MAJOR HYPERSCALERS AND CSP'S

Cloud Solution Provider

 Alibaba Cloud

 Azure

 ByteDance

 G-CORE LABS

 DIGITAL REALTY

 Cirrascale

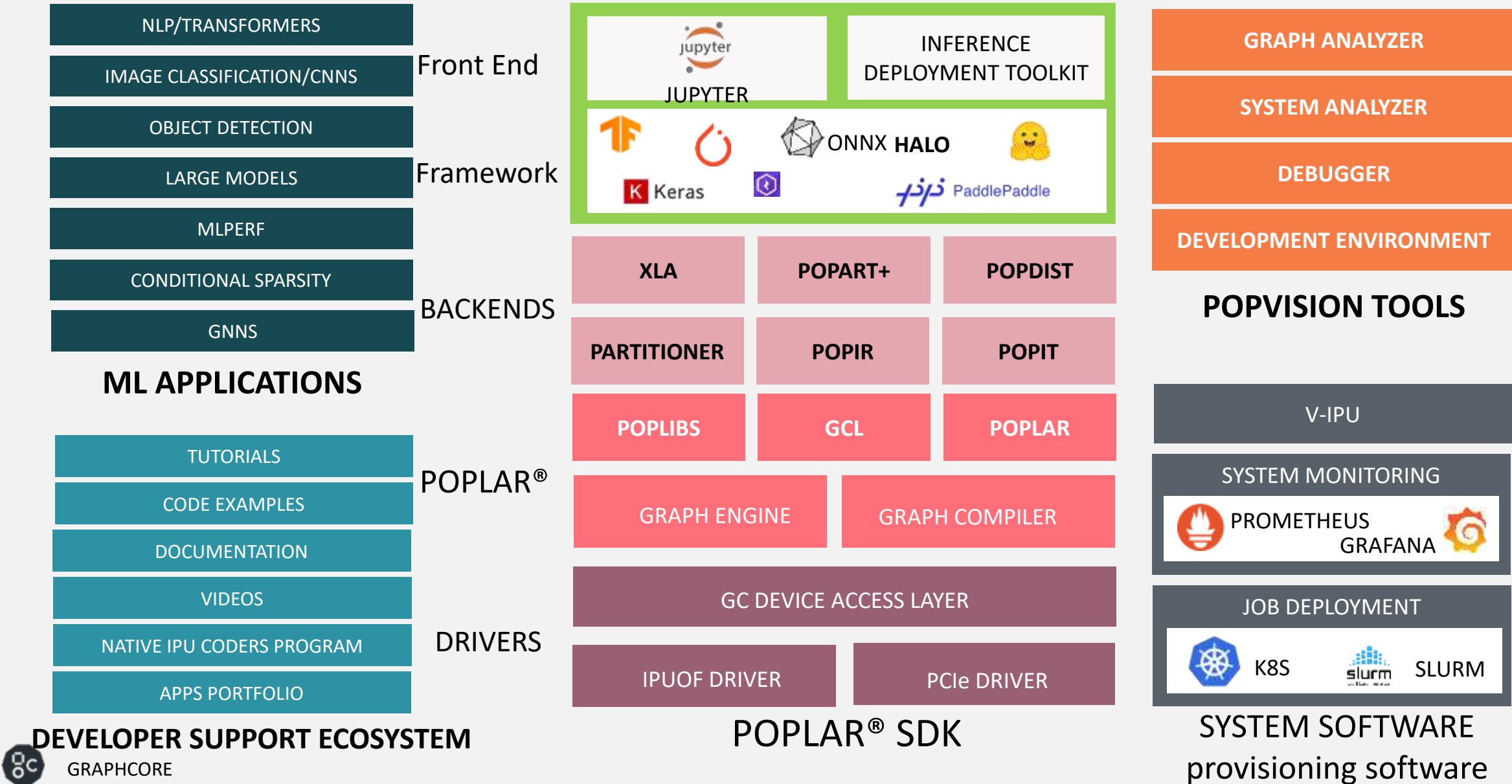


 kt

 MEGAZONE
CLOUD

Cloud Solution Provider

GRAPHCORE SOFTWARE Environment



GRAPHCORE IPUの適用領域

データセンター インターネット業界



- ・リコmend
- ・検索エンジン
- ・トレンド分析・予測
- ・スマートアシスタント
- ・画像検索
- ・顧客サポート

自動車/電機 製造業



- ・自動運転
- ・自動設計
- ・不良個所の画像認識
- ・販売支援
- ・混雑予想
- ・高度ナビゲーションシステム

ヘルスケア ライフサイエンス



- ・ゲノム研究
- ・医療画像診断
- ・疾患診断
- ・タンパク質の機能解析
- ・創薬

金融/保険/株式



- ・株価予測
- ・リスク分析
- ・景気予測
- ・マネロンの検知/防止
- ・保険料率予測
- ・顧客の本人確認・認証

研究開発 高等教育



- ・ニューラルネットによるシミュレーション
- ・ビッグデータ解析
- ・実験画像解析
- ・素粒子物理学
- ・天文学

CUSTOMER ANNOUNCEMENTS 2021



Academic Program
日本からも1団体参加



Man Group - FinTech



Tractable – Insurance AI as a service



Weather/Climate at ECMWF



Carmot Capital Fintec



NHN launch EC/IPU cloud

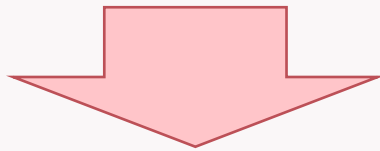


国内導入例 大手製造業、テレコム、金融、国立研究所、国立大学 など

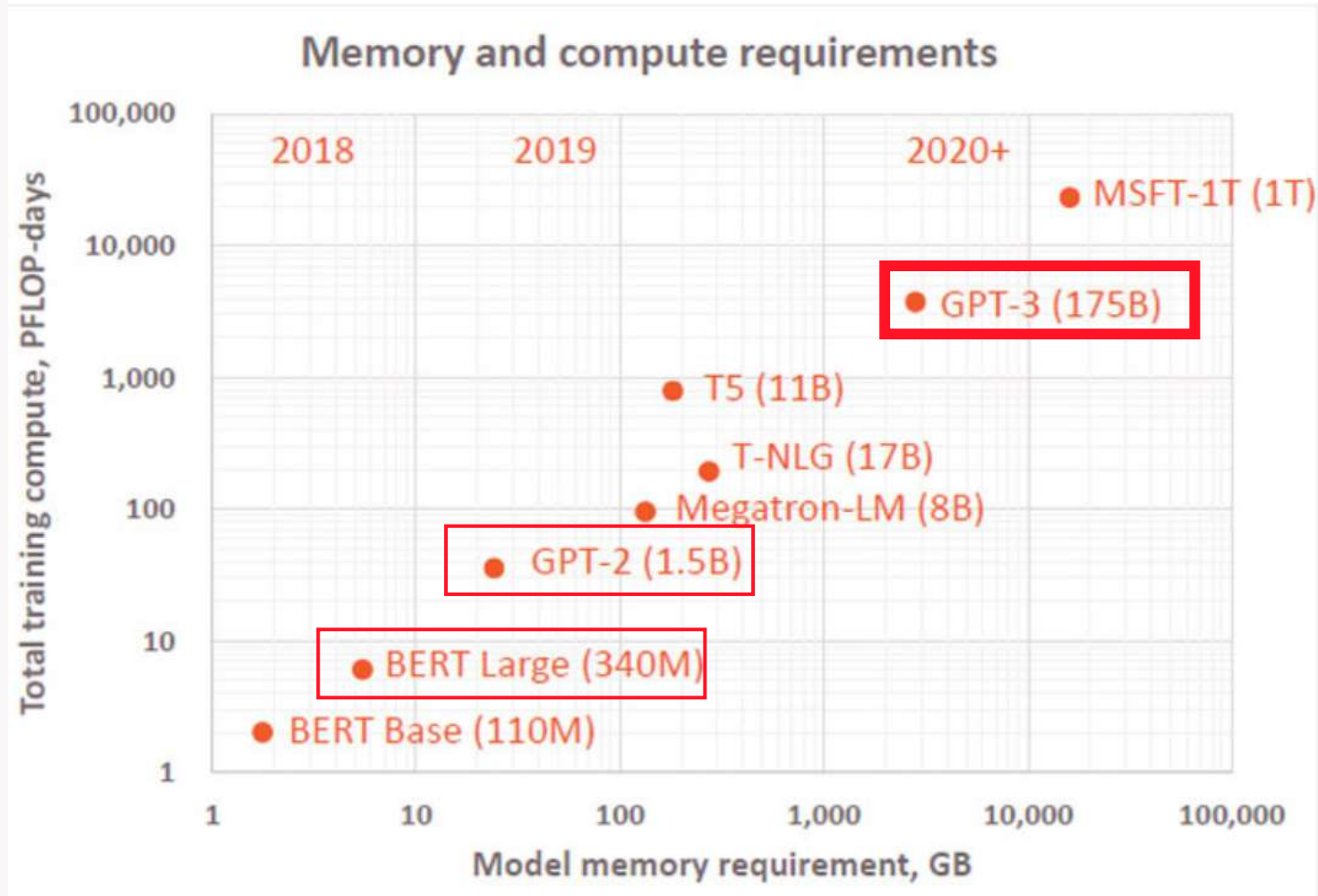
LARGE MODELへの取り組み

指数関数的に増大

- 大規模モデルの規模と計算量は2年間で約1000倍に増大
- GPT-3は1024GPUを使用し学習に2-4か月かかると言われている。
また、別の推測では学習にはV100で4096GPUで1か月、ASP利用ではコストが\$33M。
また、別ソースの推測ではA100で2000GPUで1か月、ASPでコスト\$50M)

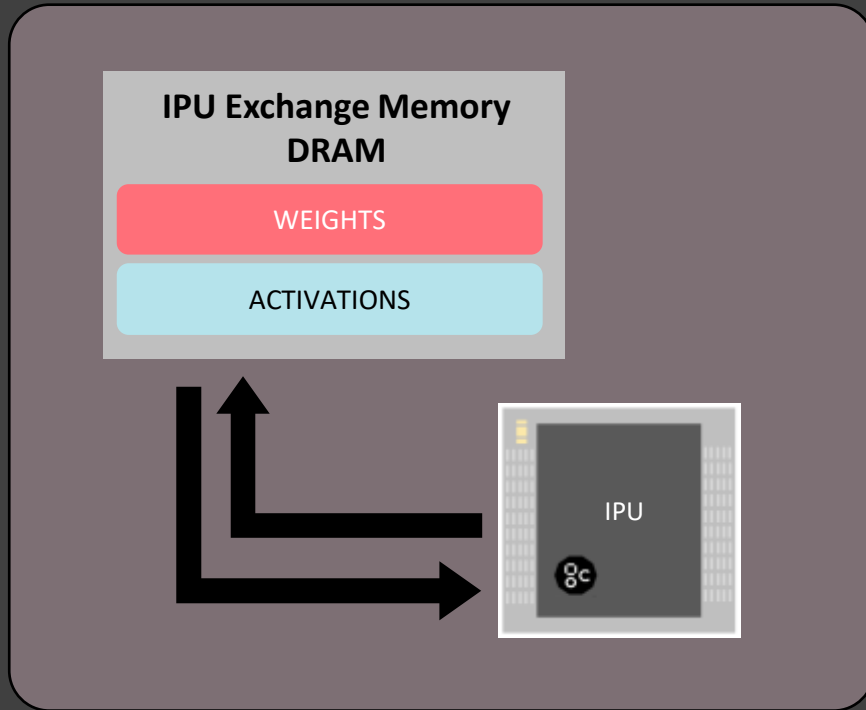


1ラック(25KW)のシステムで1日で学習可能にするには？

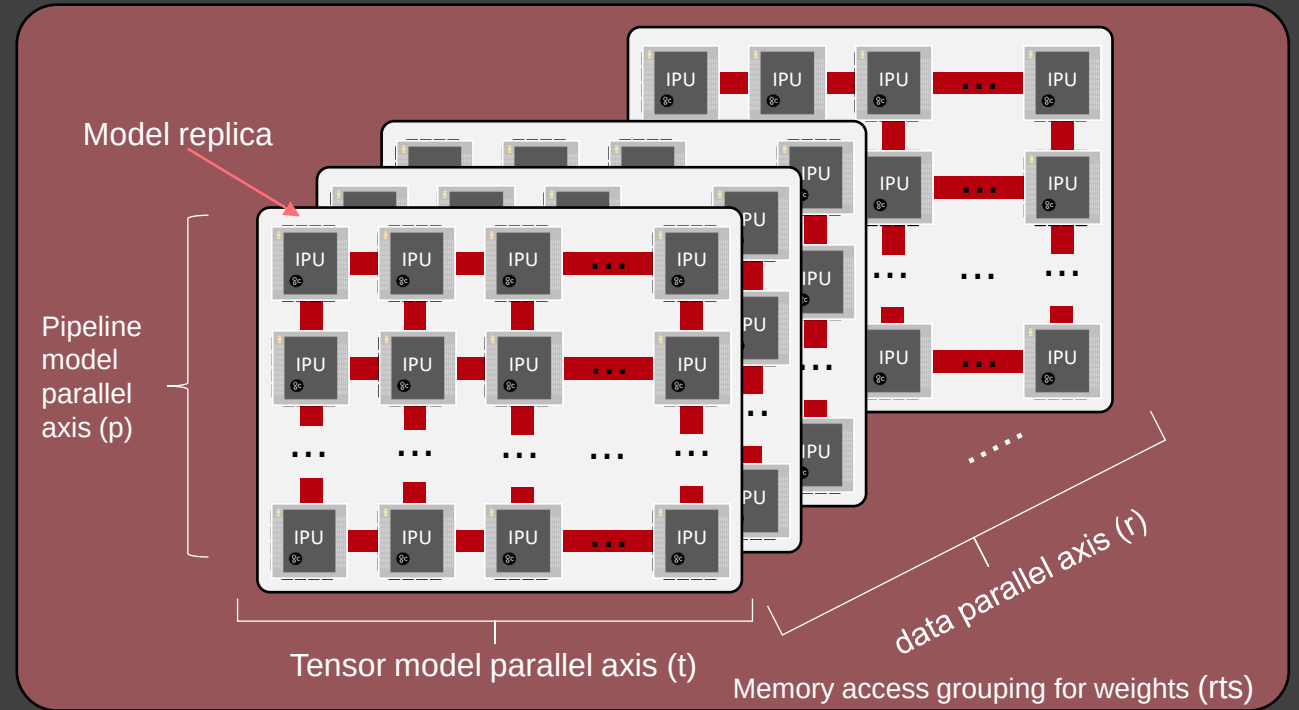


LARGE MODELへの対応

IPU-ExchangeメモリとIPU-Links™を組み合わせたIPU-PODでの実装

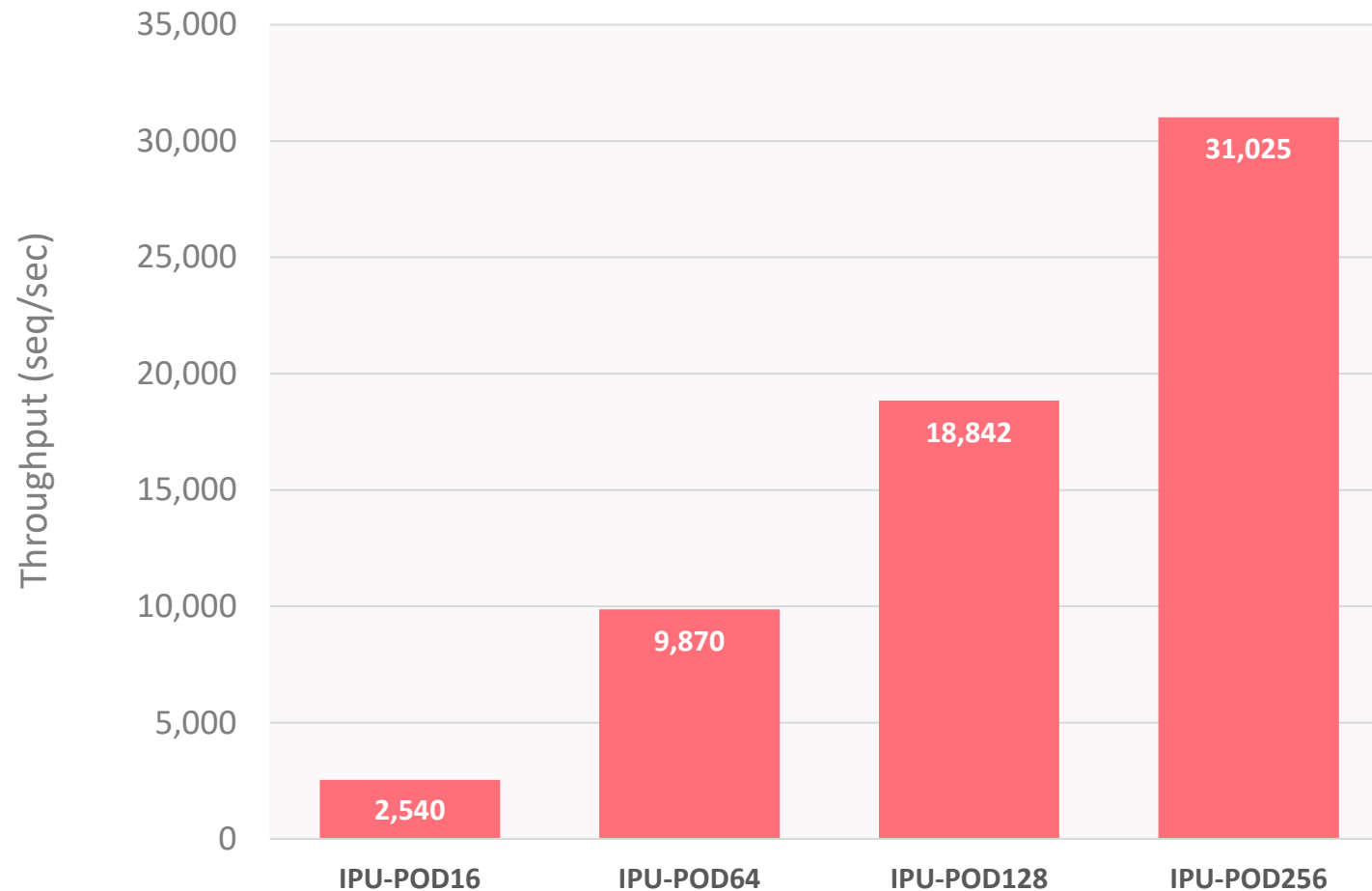


Exchange-MemoryとIn-Processor Memoryの間でウェイトとアクティベーションを“スマート（効率的に）”にスワッピングを行う。



IPU-POD全体でモデルの並列分解、Poplar®は自動化されたモデル分割をサポートします。

GPT-2 MEDIUM SCALING



NOTES: (updated 20th Dec 2021)

GPT2-Medium Training Throughput | Wikipedia dataset

IPU-POD Platforms | FP 16.16 | SDK 2.3.0 | PyTorch | <https://www.graphcore.ai/performance-results>

SUMMARY

MODEL NAME:

GPT2-MEDIUM

ML DOMAIN:

NLP

MODEL DESCRIPTION:

GPT IS A GENERATIVE PRE-TRAINING MODE.
IT IS UNIDIRECTIONAL TRANSFORMER-
BASED MODEL USED FOR PREDICTING
LANGUAGE.

DEPLOYMENT:

TRAINING

FRAMEWORK:

PYTORCH

IPU ADVANTAGE:

RESULTS ACROSS MULTIPLE IPU-POD
PLATFORMS WITH EXCELLENT SCALING



IPU-POD 性能機能評価

□ 性能機能評価メニュー

1. Graphcloudを利用した評価
2. 性能評価課題を預かり、弊社にて性能評価、報告
3. 評価用IPU-PODシステムの貸し出し
4. Try & Buy プログラム

1. Graphcloudを利用した評価

IPU-POD16 をクラウド上に配置したGraphCloudで2週間無償評価が可能。

- SSH経由
- 無償期間は2週間 (POD64 利用は有償(要相談))
- 長期利用、IPU-POD64等の利用は有償にて対応可能。
 - 弊社Webもしくは代理店様経由で利用申し込み可能
 - 事前にシステム、プログラミングの速習トレーニングも用意
 - モデル, 入力データ種別を事前にお伝えいただけますと、事前準備を対応させていただきます。



2. 性能評価課題を預かり、弊社にて性能評価、報告

評価目的、実施内容、日程などをお聞きした上で実施可能な場合、課題を預かり、弊社で移植、最適、試験実施、報告を行います。

評価を実施する時間がない、大きな規模のシステムが必要、最適化なども含めたサポートが必要な場合にご利用下さい。

3. 評価用IPU-PODシステムの貸し出し

評価目的、実施内容、日程などをお聞きした上で実施可能な場合、ホストサーバも含めたPOD-4（もしくはPOD8）を貸し出し評価。セキュリティ高いデータを外部に出せない等の受験科デモ、貸し出しシステムを利用して評価可能になります。

貸し出しに際して、

- 事前にシステム、プログラミングの速習トレーニングも用意
- モデル, 入力データ種別を事前にお伝えいただけますと、事前準備等を対応させていただきます。

4. Try & Buy プログラム

購入を前提として

評価目的、実施内容、日程などをお聞きした上で実施可能な場合、ホストサーバも含めたIPU-PODシステムをお貸出しし、評価終了後には購入いただくプログラムです。

IPU-POD16やPOD64を実際の設置環境に置き、既存システムとのインテグレーションなども含めて動作検証などを行えます。

ご清聴ありがとうございました。

Try IPU GraphCloud

性能/機能評価(POD16/2Week無償利用)
クラウドサービス(POD16,POD64)

<https://www.graphcore.ai/partners>



WebinarやVideo,のご紹介

<https://www.graphcore.ai/resources>

Graphcoreへのお問い合わせ

<https://www.graphcore.ai/>

Mail mamorun@graphcore.ai

TEL 03-3507-5600

Graphcore:dellsales@graphcore.ai



IPU-POD₆₄

Dell R6525 (3)
IPU-M2000 (16)
PowerSwitch
Rack (1) & PDU (2)