

超入門 INTRODUCTION

GPUクラスター構築

生成AI/LLM向けマルチノードGPUシステム

アーキテクチャーを理解する



GPU CLUSTER

GPUクラスター構築 超入門

生成AI/LLM向けマルチノードGPUシステムアーキテクチャーを理解する

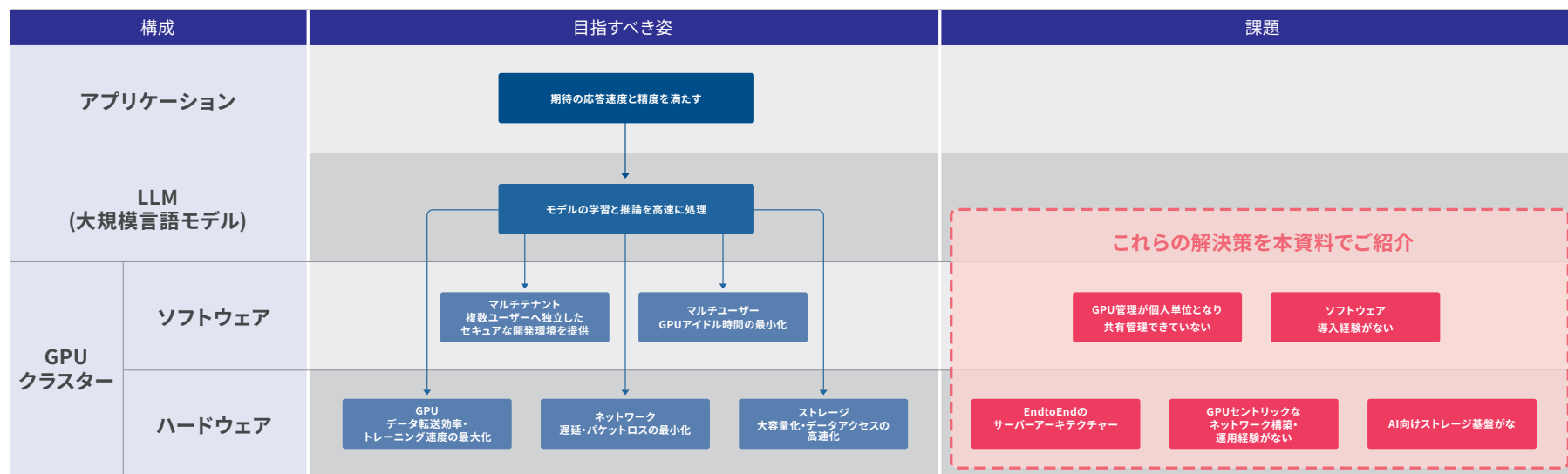
GPUクラスター 目指すべき姿と課題	2
GPUクラスター システムアーキテクチャー概要	3
各システムブロックの深掘り	4
1. GPU データ転送効率・トレーニング速度の最大化	4
2. ネットワーク 遅延・パケットロスの最小化	5
3. ストレージ 大容量化・データアクセスの高速化	6
4. マルチテナント 複数ユーザへ独立したセキュアな開発環境を提供	7
5. マルチユーザー GPUアイドル時間の最小化	8
マクニカ プロフェッショナル サポート	10
GPUクラスターに最適なソリューション	11
問い合わせ	12

LLM(大規模言語モデル)やAIモデル開発の機運が高まる昨今、GPUサーバーを複数台利用するマルチノードGPUクラスター構築の需要が増えています。

GPUクラスターは、モデルの学習と推論を高速に処理し、期待するアプリケーションの応答速度を満たす必要があります。

これらを実現するためにはAIに最適化されたGPU、ネットワーク、ストレージの環境を構築しますが、クラウド/データセンター事業者や企業のIT部門担当者にとってこれまで経験のない技術的ノウハウを求められ、課題を感じています。

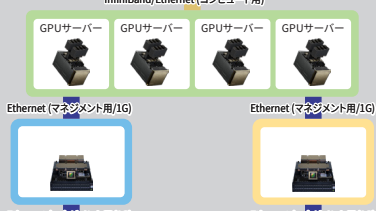
これらの課題を解決する糸口として、AIワークロードに最適化されたGPUクラスター構築の設計ポイントを本資料で紹介します。



まずは理想的なGPUクラスターの全体像を紹介します。

GPUクラスターではGPU・ネットワーク・ストレージのインフラ構築だけでも大変な作業ですが、ユーザー利用を考慮したソフトウェア構築までを設計、実装をして初めてAIに最適化されたGPUクラスターが完成します。

ここからは実際の構築順序に沿って各システムブロックを深掘りしていきます。

目指すべき構成	目指すべき姿	実現する手段	構築順序
	マルチテナント 複数ユーザーへ独立した セキュアな開発環境を提供	マルチテナント対応 最適なアプリケーション実行環境	4
	マルチユーザー GPUアイドル時間の最小化	ジョブスケジューリング導入	5
	ネットワーク 遅延・パケットロスの最小化	InfiniBand fabric / Ethernet BGP fabric選択 RDMA/RoCEv2接続 Adaptive Routing採用	2
	GPU データ転送効率・ トレーニング速度の最大化	Rail-Optimized Topology 採用 GPU-Direct RDMA採用	1
	ストレージ 大容量化・ データアクセスの高速化	チェックポイント対応 マルチプロトコル対応	3

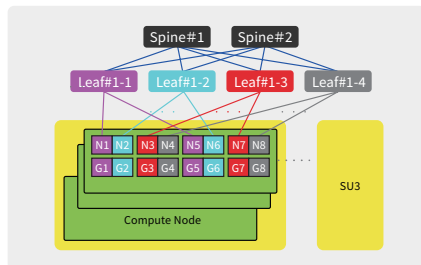
課題 生成AIデータセンターのネットワーク構築

- マルチテナント要件 (Ethernet) を満たしながら、密結合プロセス・低ジッター・高バースト性トラフィックへの対応を両立すること
- サーバーアーキテクチャー変化への対応
(ポート数増加、広帯域化、PCIeのボトルネック化)

Control / User Access Network (N-S)	AI Fabric (E-W)
疎結合アプリケーション	密結合プロセス
TCP (狭帯域フロー)	RDMA (広帯域フロー)
High Jitter Tolerance	Low Jitter Tolerance
Oversubscribed Topologies	Nonblocking Topologies
多種多様なトラフィック、マルチパス	Burst型トラフィックのピークパフォーマンス対応

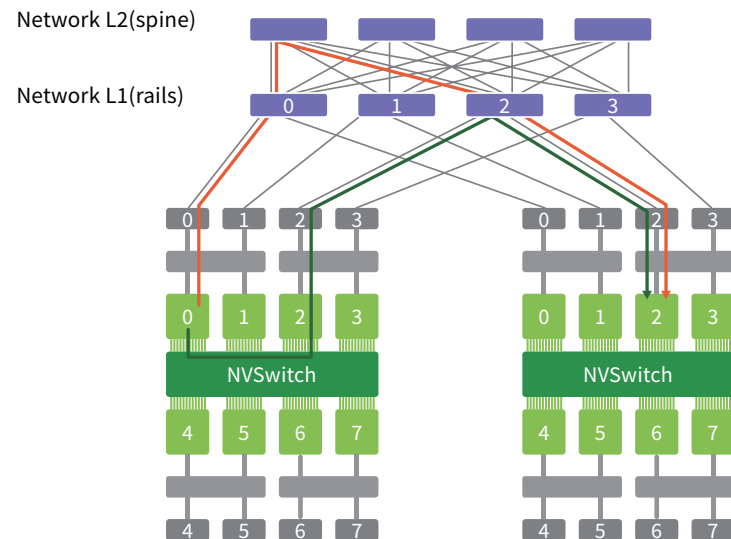
対応策 Rail-Optimized Topology

- GPU番号ごとに接続する
Leafスイッチを分散
- 同ノード内GPUは専用経路にて接続
(NVLink, NVSwitch)



詳細 Rail-Optimized Topology のメリット

- NVLink/NVSwitch により、全ノードの同一番号GPUに1Hopで到達
→ネットワーク遅延の最小化/均一化、Spineへのトラフィック集中低減
- NCCL (NVIDIA Collective Communication Library) における経路選択によりEthernet/NVLink/PCIe/QPI全ての経路を把握
→最速・最適な経路を選択し、広帯域・低遅延な通信を実現



課題 GPUセントリックなネットワーク構築・運用

- 複数のユーザーへ独立かつセキュアな開発環境の提供
- GPUリソースの最大効率化、帯域幅確保
- データ転送非効率及びネットワークのコンジェスション
- 最適化されていないルーティング
- 複雑なプロトコルの併用

対応策 アクセラレーション技術を用いた効率化

各種技術と適応するネットワークの種類

Compute Network

- InfiniBand/Ethernetファブリックの選択
- RDMAやRoCEv2
- GPU-Direct RDMA
- Adaptive Routing

Storage/In-Band Management Network

- InfiniBand/Ethernetファブリックの選択
- RDMA/RoCEv2

詳細 ネットワークファブリックの構成要素**Compute Network**

GPUサーバー間の高速な通信を実現するネットワーク

Storage Network

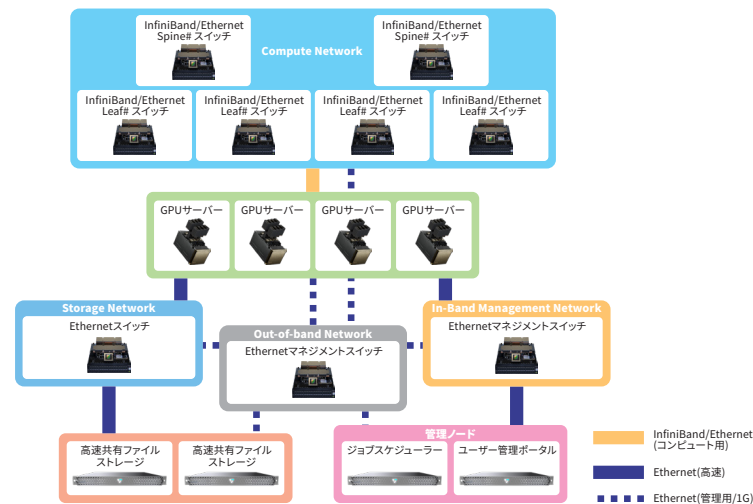
ストレージ GPUサーバー間の大容量かつ高速アクセス可能なネットワーク

In-Band Management Network

ユーザーアクセス、ジョブ投入等のクラスター内サービス用ネットワーク

Out-of-Band Network

各機器のOS、BMC管理用のネットワーク



課題 AI向けストレージ基盤構築

- データパイプラインの各レイヤーサイロ化を撤廃
- スループット・容量のスケラビリティを検討
- データのバックアップ(チェックポイント)の考慮
- クラウド利用への対応のためマルチテナント要件を満たす

対応策 マルチプロトコル・チェックポイント対応

- マルチプロトコル対応可能
- スケーラブルなアーキテクチャの採用
- 運用が容易で堅牢なバックアップ手法のあるストレージの採用
- マルチテナント対応可能

詳細

マルチプロトコル対応がなぜ必要か？

データパイプラインにおいて、ストレージプロトコルは複数必要となるため、マルチプロトコル対応のストレージが必要となる。

例) 学習(NFS)、Data lake(S3)、周辺ツール/DB(Block)

チェックポイントの重要性

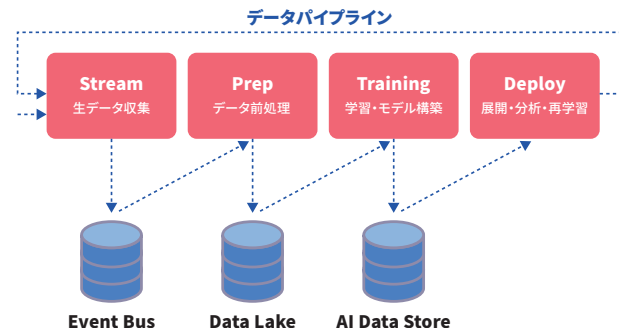
ストレージのデータは、高速アクセスの目的でメモリー内に保管される。一般的にストレージはチェックポイントを使用してメモリー内のデータをディスクにコピーすることによって、データ永続性をサポートしている。これにより、計画的なシャットダウンや、または予期しないシャットダウンが発生した後にデータをリカバリーすることができる。

AIワークロードのための高速処理

必要なデータを迅速に取得できれば、データ処理や分析の効率が向上する。そのため、ファイルのアクセス速度を向上するためのプロトコルやソフトウェアが対応していることが重要となる。

マルチテナント対応

AIクラウドで採用されるストレージでは、マルチユーザー毎のストレージアクセスができることが要求されるため、ストレージ自体もマルチテナント対応していることが重要となる。



課題 複数ユーザーでGPUリソースの共有

- ・膨大なGPU資源で構成されるGPUクラスターは、複数のユーザーで分配して使用される
- ・複数社のユーザーが混在する場合に、会社間/部門間でのセキュリティの担保と同時に、計算資源の最適化を求められる

対応策 マルチテナント対応

<マルチテナント対応の目的>

- ・複数テナント間をセキュアに分割し、且つ計算リソースを無駄なく使用できる

<マルチテナント対応のポイント>

- ・ネットワーク/ストレージなどのインフラ機器、ジョブスケジューラーなどのソフトウェア、Gateway/Firewallなどの認証機器、様々な視点での対応が必要

詳細 マルチテナント化で検討すべき事項

- ①データセンターへのアクセス管理
テナント単位、ユーザー単位、プロジェクト単位の認証/アクセス管理
- ②ジョブスケジューラーによるユーザー管理/リソース管理
各単位毎のログイン管理/プロジェクト毎のログイン管理
各単位毎の計算リソース管理/最適化
- ③ストレージアクセス管理
各単位毎にアクセス管理/認証
各単位毎にファイルシステムの分離/暗号化
ファイルアクセス権限の管理
ストレージリソース最適化
- ④ネットワーク管理
テナント毎のネットワーク分離
QoSによるノイジーネイバー対策(ネットワークリソース最適化)

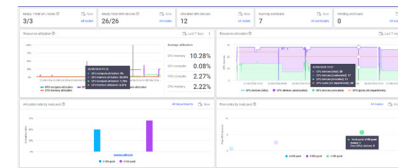


課題 GPUを共有管理/効率利用

- ユーザー管理/アクセス制御の仕組みを最適化
- GPU、CPU、メモリーリソースの共有化/最適化
- アイドル状態のGPUが存在するのかどうかの把握
- アイドル状態GPUの有効活用

**対応策 ジョブスケジューラーによる
ユーザー管理/リソース管理**

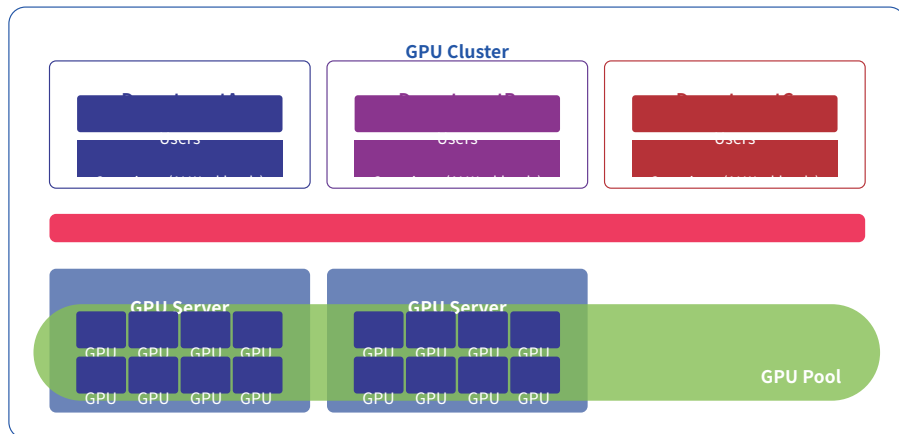
- ユーザー管理とアクセス制御の設定を一元的に実行できるジョブスケジューラーの導入
- リソースをプール化して柔軟にマルチユーザーへリソース共有できる仕組みを導入
※両方を実現できるソフトウェアの導入が望ましい
- 効率的かつ柔軟にGPUリソースを利用可能にし、アイドル状態のGPU有効活用を実現
- 付随する監視機能でアイドル状態のGPU有無を定期的に監視しながらGPUを管理



詳細

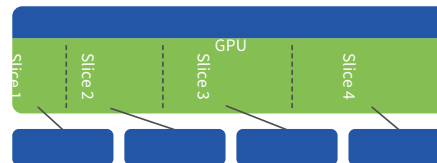
ユーザー管理とアクセス制限/
マルチユーザへリソース共有

- 同一Cluster内でDepartmentという単位で領域を分けることが可能
- ユーザー、WorkloadをDepartment単位で管理することでよりセキュアにクラスターを運用可能
- GPUリソースの共有使用で効率利用の実現
- ユーザーはGPUと合わせて最適なコンテナ環境を利用可能



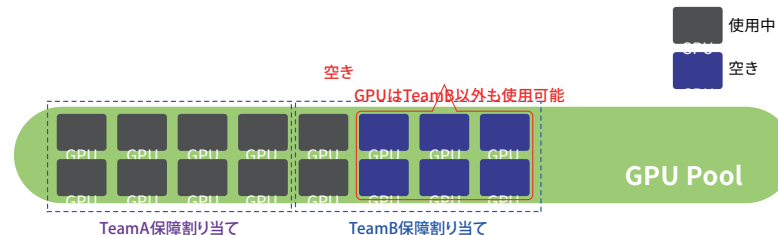
GPUリソースの分割

- GPUメモリをSliceして複数コンテナでGPUのシェアが可能
- 各Slice領域は互いにIsolateされているため安全に使用可能



動的なGPU割り当てポリシー

- ユーザー (またはチーム)が確実に使用できるGPU割り当て(保障割り当て)を実施したうえで空きGPUが存在する場合は割り当て数を超えてGPUリソースを活用可能
- 保障割り当て範囲でGPU利用が最優先される



マクニカプロフェッショナルサポートは、GPUクラスターに最適な最新NVIDIAハードウェア(GPU、ネットワークシステム)と高速化・効率化に必要なソフトウェアを選定し、弊社スペシャリストによる高度なインテグレーションを提供します。
要件定義からインテグレーション、導入後まで伴走し、お客様へ最高クラスのAI開発環境をお届けします。

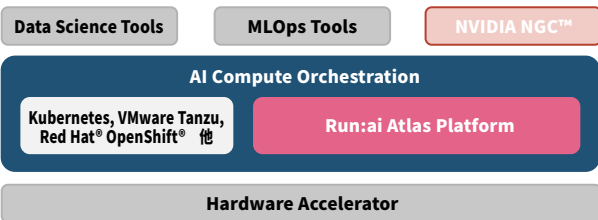
サポートフロー





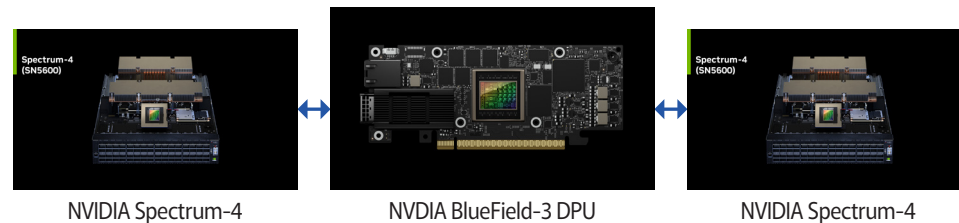
NVIDIA DGX™ H200/B200/B300 Supermicro社 NVIDIA HGX H200/H300

NVIDIA DGX SuperPOD™ の基盤で、画期的なNVIDIA B300 TensorコアGPUを搭載し、AIの活用を推し進めます。



Run:ai Atlas Platform

GPUリソースの動的かつ自動的なチーム間共有を実現するプラットフォームで、Kubernetesのプラグイン(ソフトウェア)として提供されます。



NVIDIA Spectrum-X

Ethernetスイッチ「NVIDIA Spectrum-4」とネットワークアクセラレーター「NVIDIA BlueField-3 DPU」で構成され、1.7倍のAIパフォーマンスと電力効率を実現します。



VAST Data Platform

AIインフラにおけるストレージの課題を解決するために、次世代のデータプラットフォームを提供しています。
マルチプロトコル (NFS, S3, CIFS, etc)に対応し、すべてのデータを統合管理(Structured, Unstructured)できます。



・本資料に記載されている会社名、商品、サービス名等は各社の登録商標または商標です。なお、本資料中では、「™」、「®」は明記しておりません。・本資料は、出典元が記載されている資料、画像等を除き、弊社が著作権を有しています。・著作権法上認められた「私的利用のための複製」や「引用」などの場合を除き、本資料の全部または一部について、無断で複製・転用等することを禁じます。・本資料は作成日現在における情報を元に作成されておりますが、その正確性、完全性を保証するものではありません。

お問い合わせ

株式会社マクニカ クラビス カンパニー NVIDIA製品担当

〒222-8561 横浜市港北区新横浜1-6-3 マクニカ第1ビル



045-470-9825



clvinfo@macnica.co.jp



<https://www.macnica.co.jp/business/semiconductor/support/contact/>