

GMO GPUクラウド

のシステム構築事例紹介 &
事例をもとにGPUネットワークのトレンドを
NVIDIAさんに聞いてみよう

GMO INTERNET



登壇者

GMOインターネット 友 源輝

NVIDIA

木戸 大二郎

1. 登壇者自己紹介
2. GMO GPUクラウドのシステム構築事例紹介
 - 2-1. GMOインターネットの紹介
 - 2-2. システム構成
 - 2-3. ネットワーク構成
3. NVIDIAさんに聞いてみよう
 - 3-1. NVIDIAの紹介
 - 3-2. NVIDIAさんに質問をしてみよう
 - 3-3. 会場を交えての議論

友 源輝 Genki Tomo

所属：ネットワーク技術チーム

出身：大分県

趣味：アウトドア、釣り🐟、狩猟採取🦊🍄

2007年 上京してネットワークエンジニアになる

2007年～2017年 NW運用、SIerで中小の構築に携わる

2018年 GMOインターネット入社

ConoHaなどクラウド商材のNWインフラ業務に携わる

NWインフラの設計・構築・運用まで一通り担当

2024年 GPUクラウドプロジェクトにジョイン、GPUクラウド環境のNW設計・構築を担当



木戸 大二郎

NVIDIA 合同会社

出身：北海道 千歳市、苫小牧市

趣味：PCゲーム

現職：NVIDIA ネットワーキング担当 ソリューションアーキテクト

前職：Cisco 10年在籍 (WAN Routing/ Datacenter Switching/ Cloud Networking/ AI Networking)

前前職：Cyberagent でデータセンターネットワーク設計、構築、運用

それ以前：キャリア系ネットワーク設計、構築、運用を8年ほど



GMO GPUクラウド システム構成

1. 登壇者自己紹介
2. GMO GPUクラウドの構築事例紹介
 - 2-1. GMOインターネットの紹介
 - 2-2. システム構成
 - 2-3. ネットワーク構成
3. NVIDIAさんに聞いてみよう
 - 3-1. NVIDIAの紹介
 - 3-2. NVIDIAさんに質問をしてみよう
 - 3-3. 会場を交えての議論

GMO INTERNET

主要事業

インターネットインフラ事業

ドメイン・レンタルサーバー
(ホスティング) 事業

インターネット接続
(プロバイダー) 事業

インターネット広告・メディア事業



ドメイン取るなら 公式登録サービス
お名前.com .com .jp .co.jp
<https://www.onamae.com> by GMO



GMO GPUクラウド

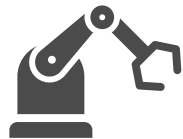
AI指示文・プロンプトなら

教えて.AI
by GMO

サービス名：GMO GPUクラウド

提供開始日：2024年11月22日

生成AIモデル学習や 機械学習に最適化されたGPUクラウドサービス
インフラチューニング不要な計算資源を提供



生成AIの急速な進化
さらなる需要拡大

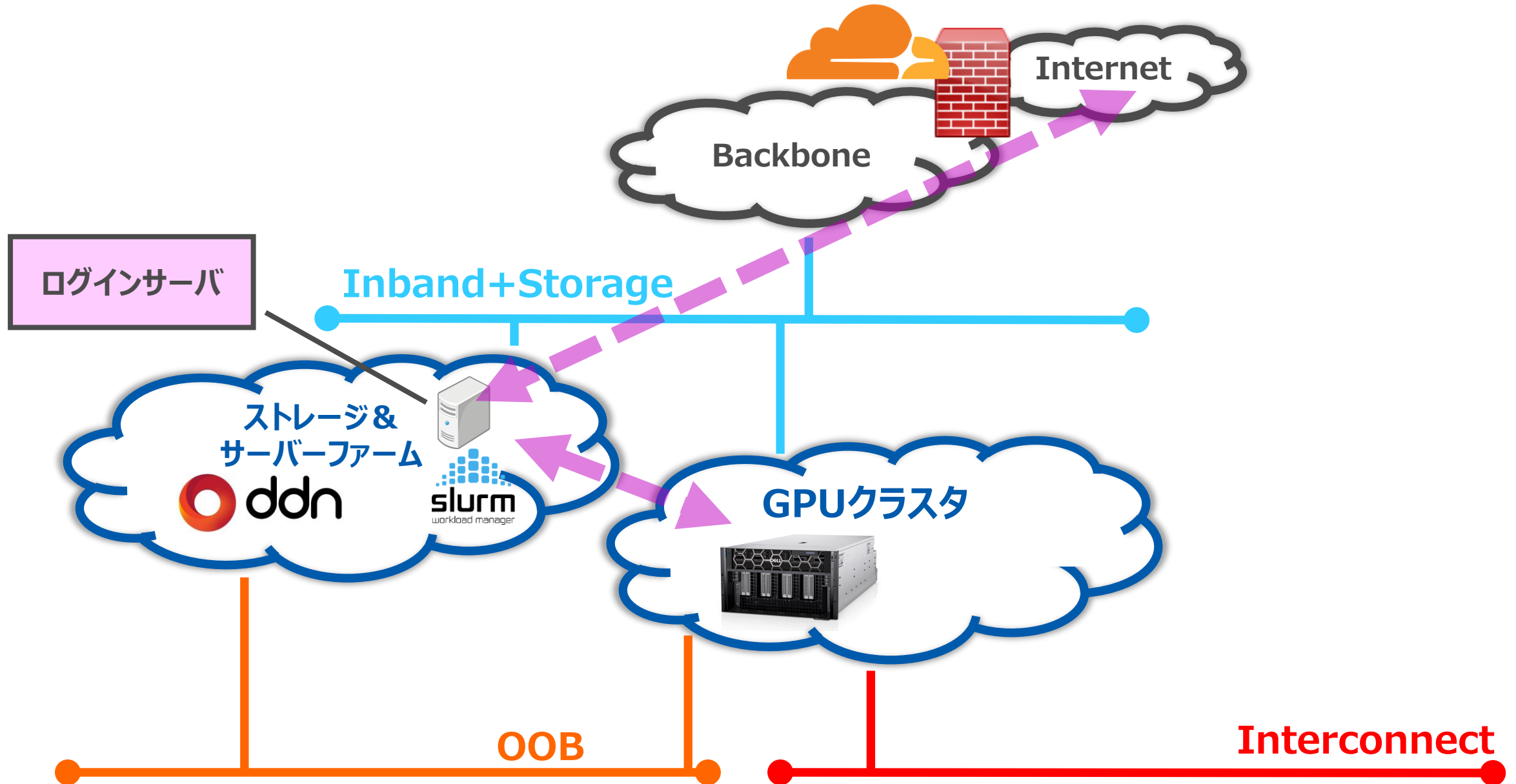


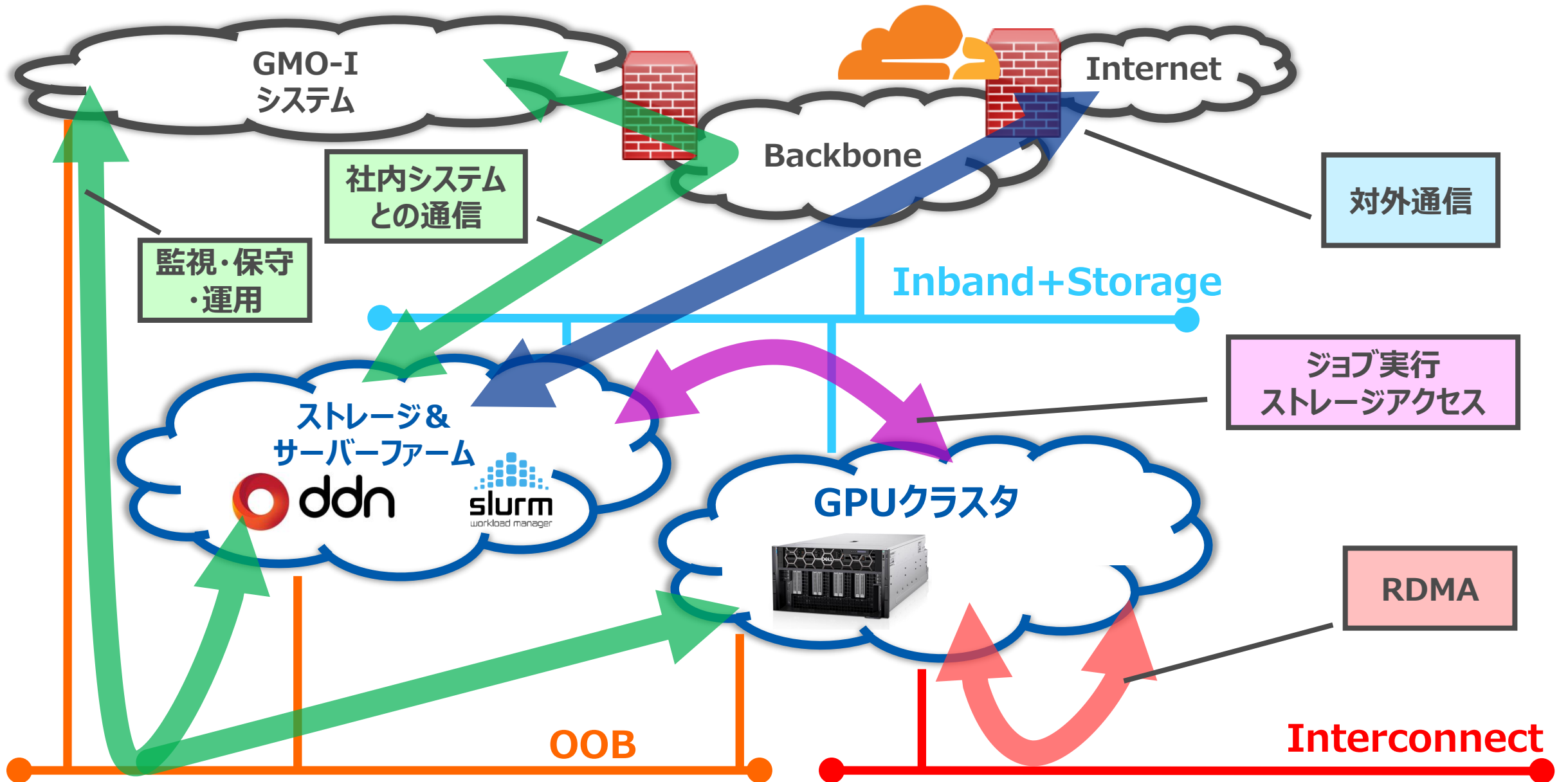
計算リソースの需要の
増加

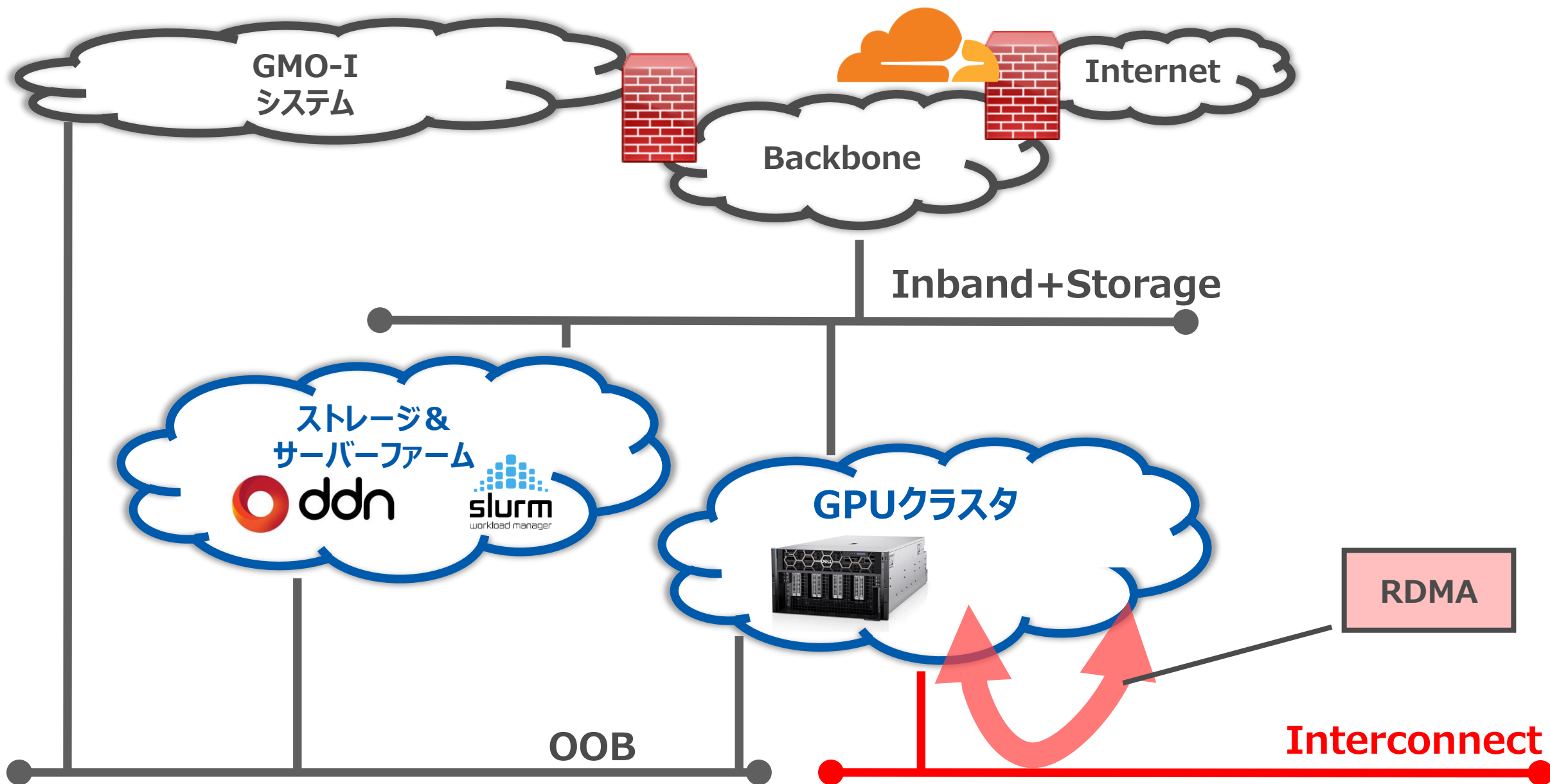


日本のAI研究開発の加速を促し
国際競争力を高めるために世界最高水準
の計算リソースの提供を目指す

システム構成







GPUノード筐体は高密度GPUサーバ
DELL Power Edge XE9680 を採用。

NVIDIA H200 GPUを 8枚搭載
NVIDIA BlueField-3 SuperNICを 8枚 IC接続



Interconnect-SWに、
NVIDIA Spectrum-X対応 SN5600を採用



Ethernet-SW
800G OSFP ×64

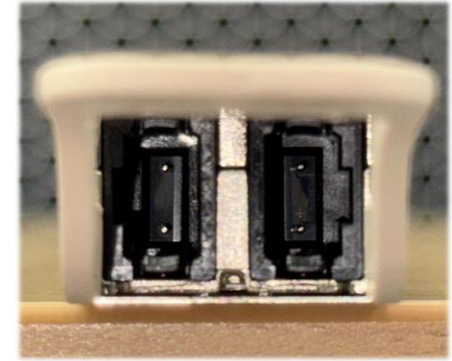
Q. イーサネットとインフィニバンド、
どっち選べばいいの？

Interconnect-SWに、
NVIDIA Spectrum-X対応 SN5600を採用

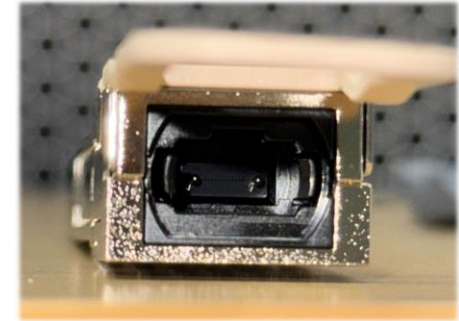


Ethernet-SW
800G OSFP ×64

Interconnect-SW側 :
OSFP 400G 2xSR4



GPUノード側 :
QSFP112 400G-SR4



ケーブル :
MPO12 マルチモード





NVIDIA Spectrum-X対応

- RoCEv2

- RDMA通信のIBパケットをカプセル化
加えて、ECN・PFC などを利用し、輻輳制御

- アダプティブブルーティング

- インターフェイスごとの通信を監視し、動的にバランシング

- NVIDIA独自の輻輳制御

- BlueField-3 SuperNICと連携し、エンドツーエンドの輻輳制御

Q. 輻輳制御って
そんなに大事なの？

NVIDIA Spectrum-X対応



- RoCEv2

- RDMA通信のIBパケットをカプセル化
加えて、ECN・PFC などを利用し、輻輳制御

- アダプティブブルーティング

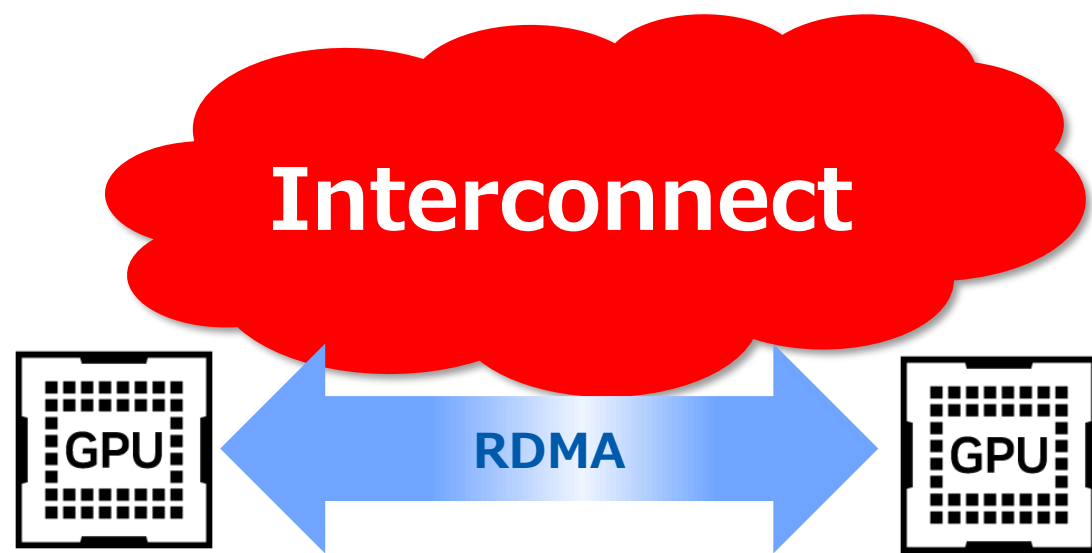
- インターフェイスごとの通信を監視し、動的にバランシング

- NVIDIA独自の輻輳制御

- BlueField-3 SuperNICと連携し、エンドツーエンドの輻輳制御

インターコネクト ネットワーク 構成

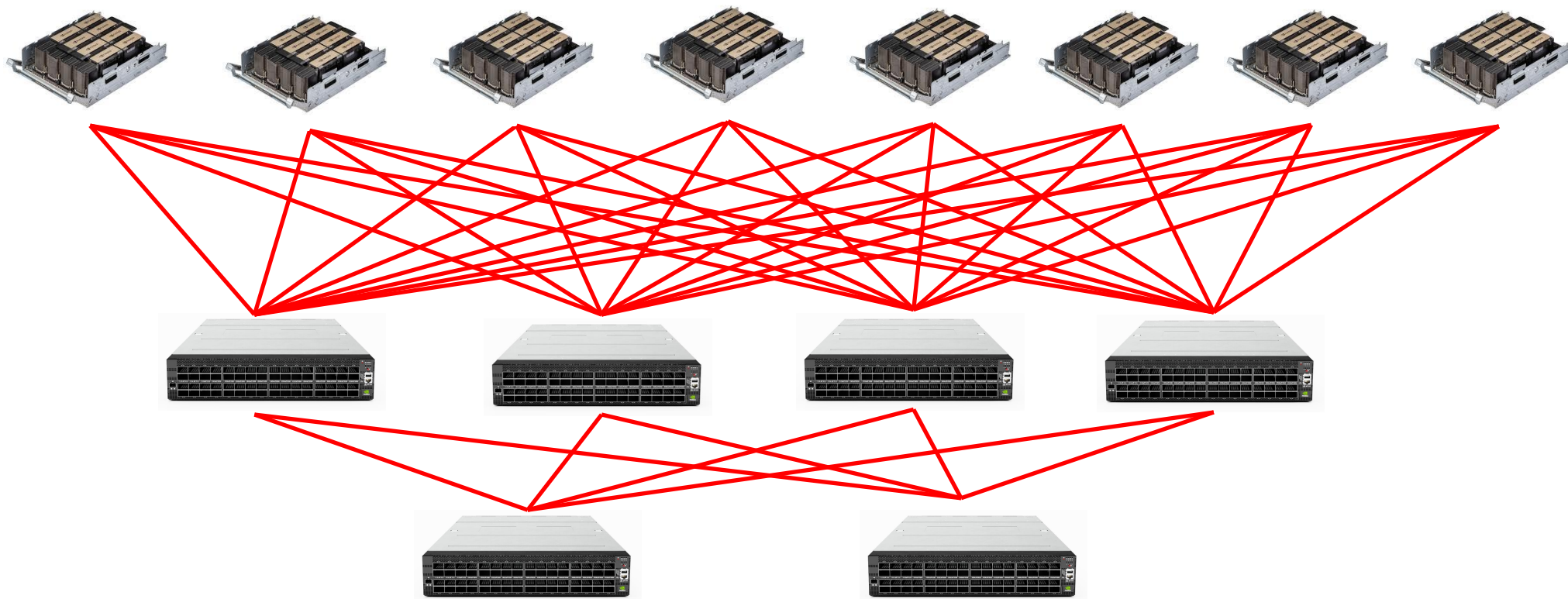
GPUネットワークの要、 Interconnectネットワーク



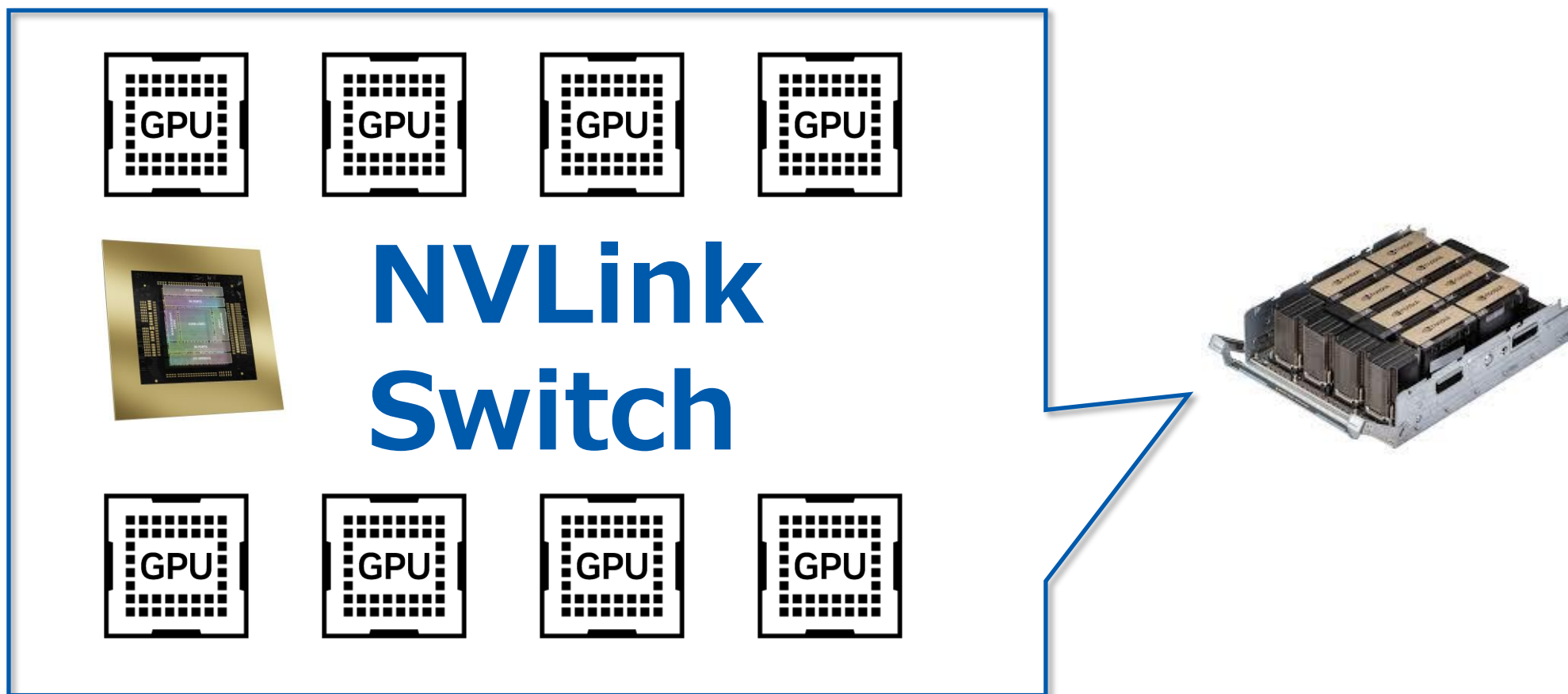
GPUノード間の**RDMA**
(Remote Direct Memory Access)
を行うネットワーク

パケットロスが計算の失敗に直結、
ロスレスNWを構築する必要がある。

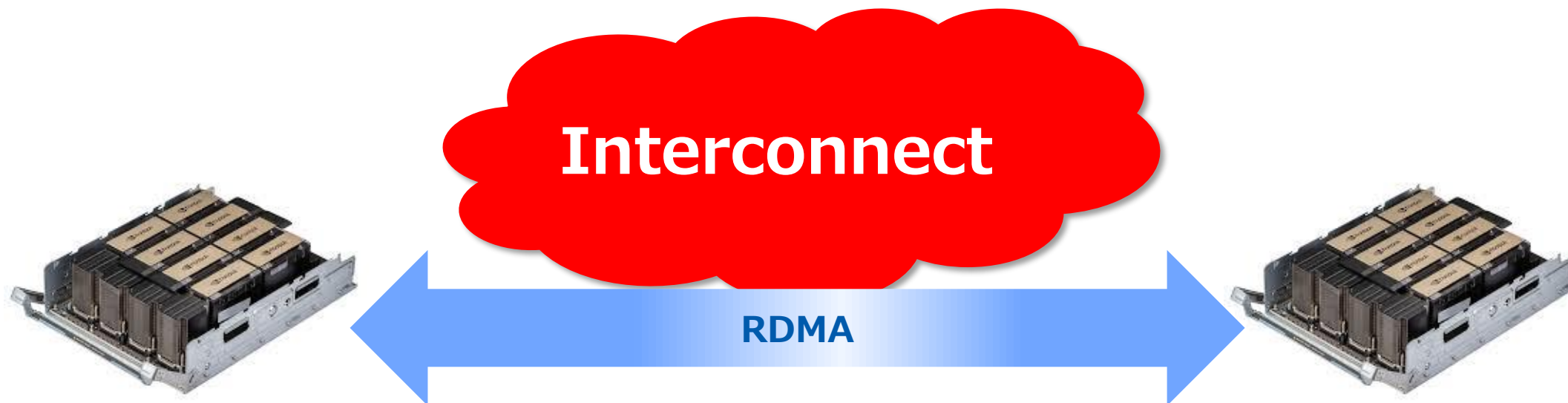
ロスレスネットワークの実現のため、
トポロジには **Rail optimized topology Full Bisection**を採用



GPUノード内のGPU間通信は、
サーバ内蔵の**NVLink Switch**で処理



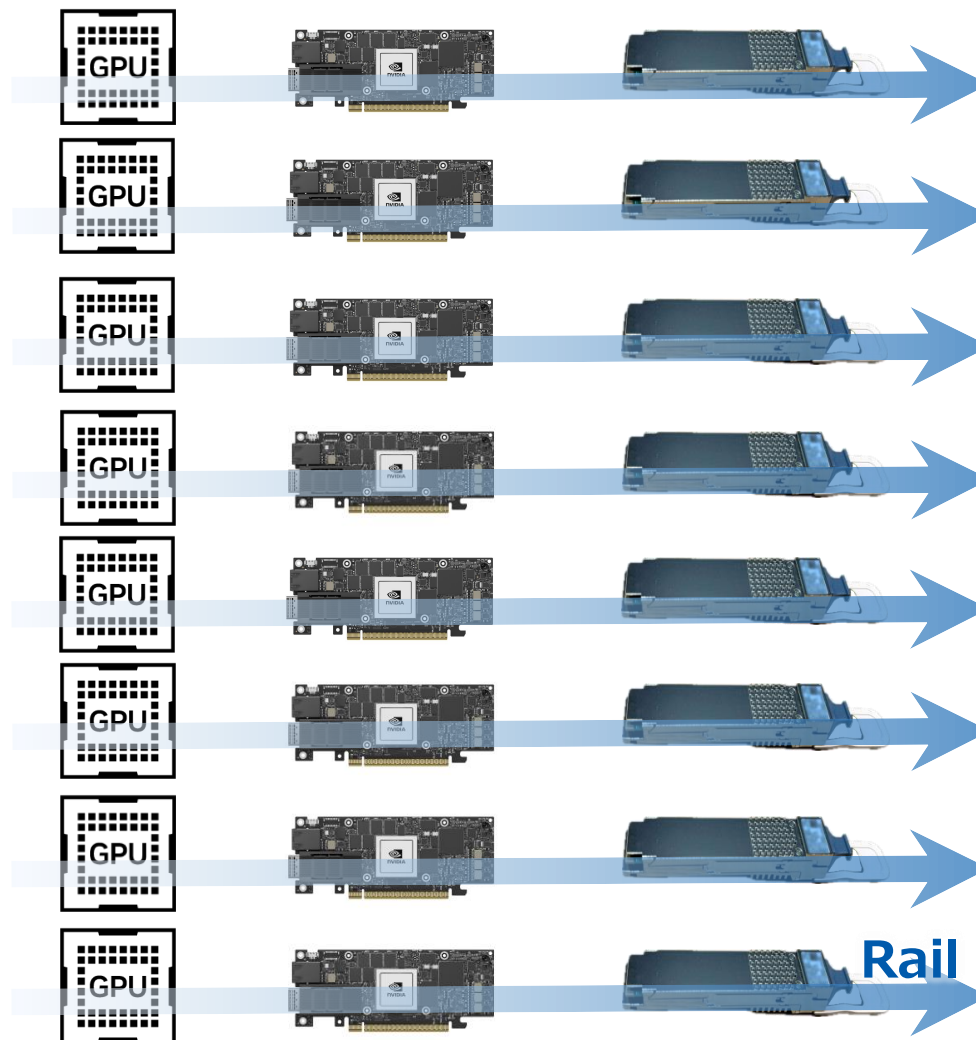
1ノードで処理しきれないほどの大きな計算は、
複数のGPUノードで連携して処理





GPUノード間通信には
GPU**1枚**につき
400Gb/s
の帯域を確保

GPU BF3NIC 400Gb-Optics



Interconnect
経由で他GPU
ノードへ

Rail

GPUノード#1



GPU#1



GPUノード#2



GPU#1



GPUノード#3



GPU#1



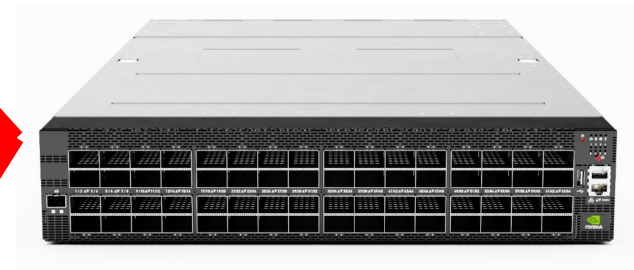
GPUノード#4



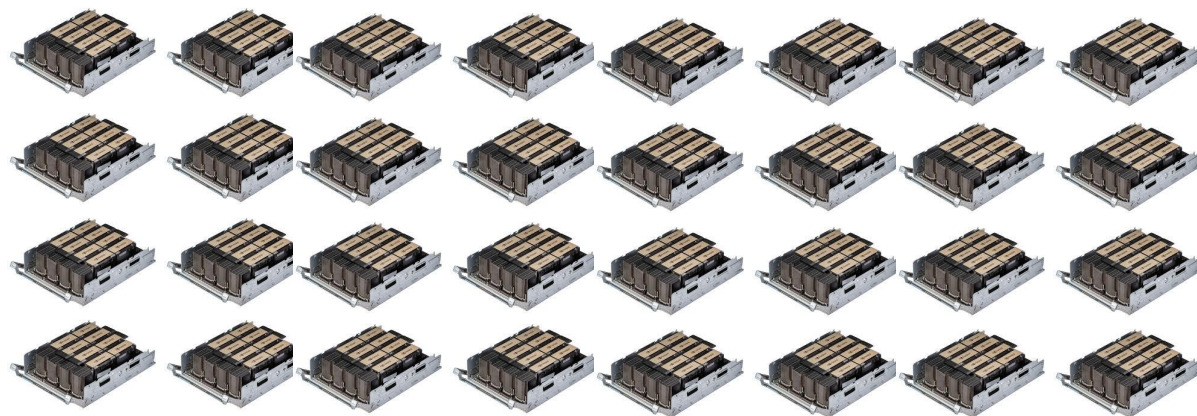
GPU#1



各GPUノードの
同じGPU番号を
同じアップリンクSWへ

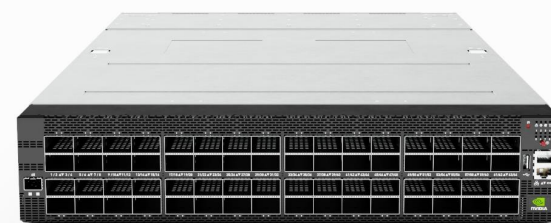
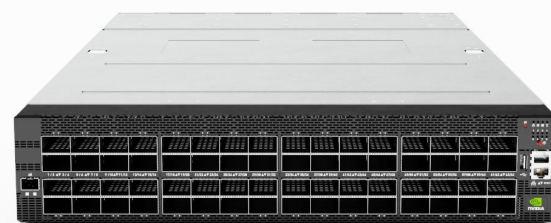
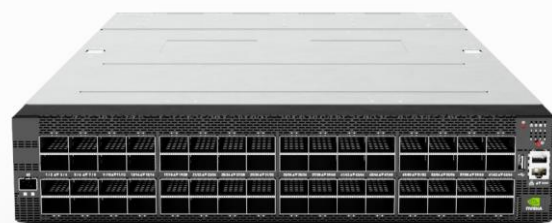
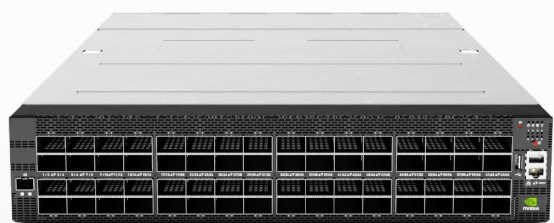


Leaf-SW



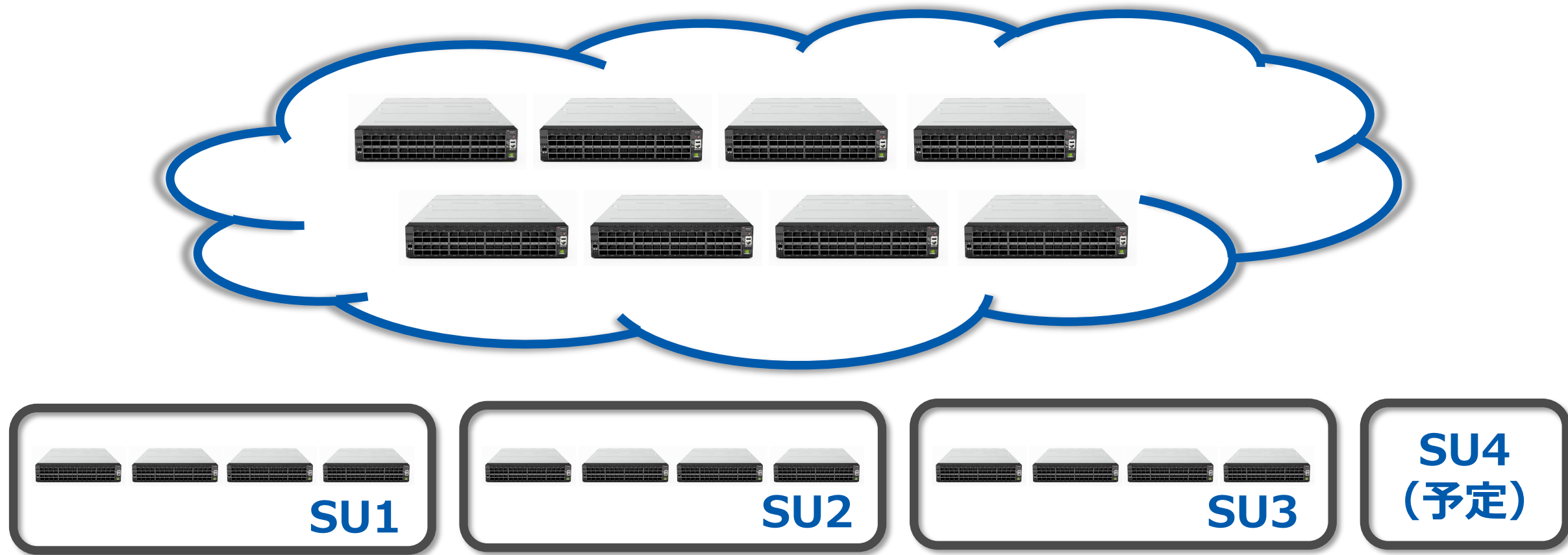
GPUノード **×32台**
=1SU(スケーラブルユニット)

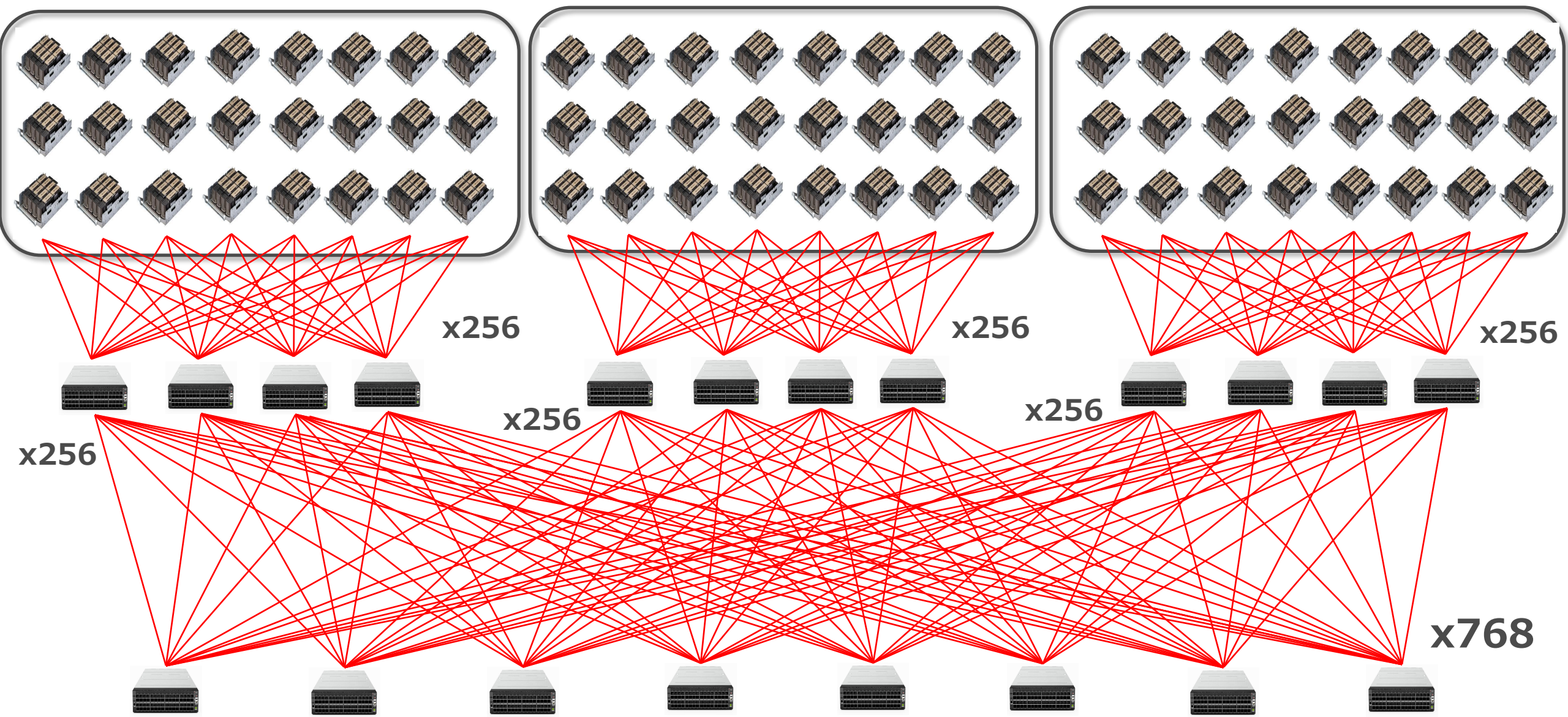
アップリンクケーブルは
 $8 \times 32 = 256$ 本



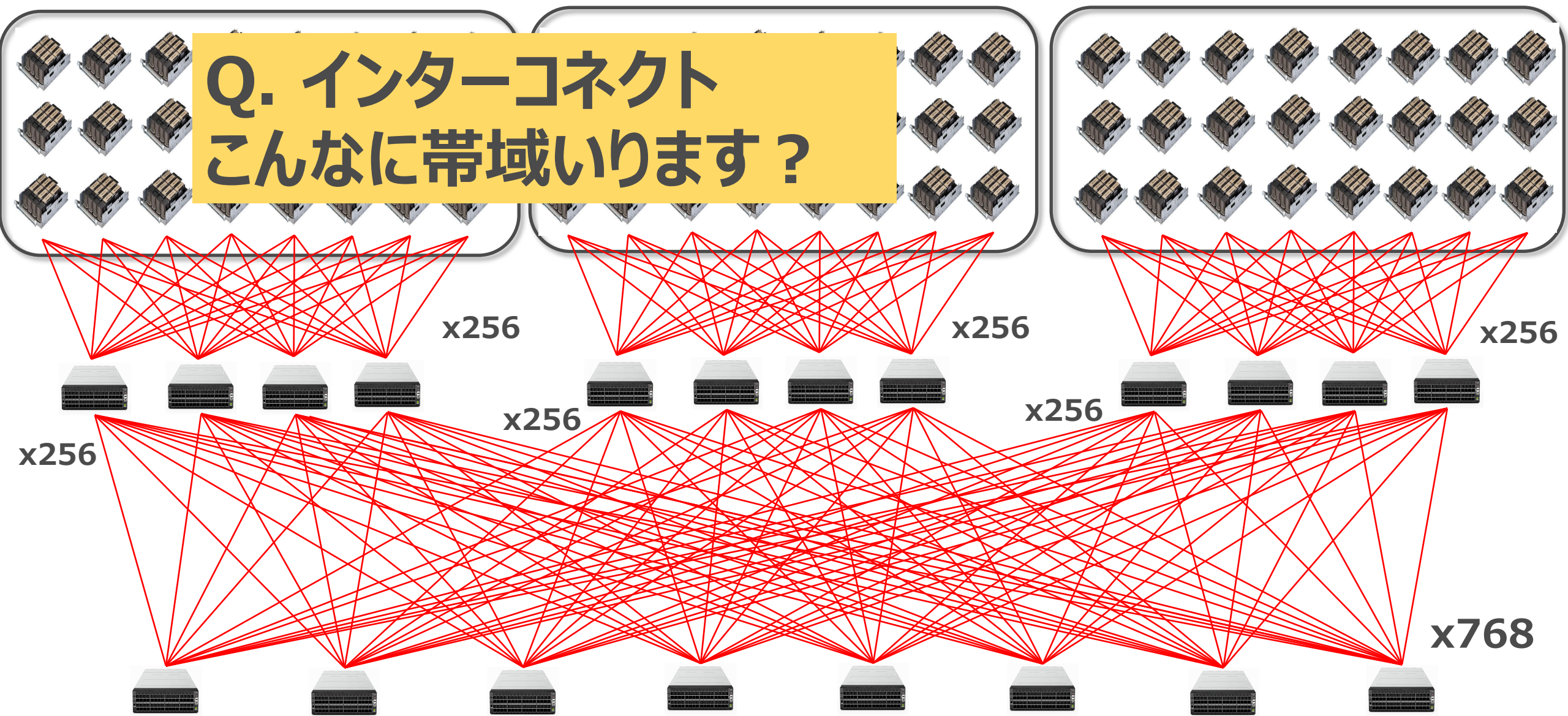
Leaf-SW **NVIDIA SN5600** x4 で収容

SU(スケーラブルユニット)間の通信を処理する
Spine-SWは **NVIDIA SN5600** x8 を配置、
4SU(leaf16台)を収容できる構成。





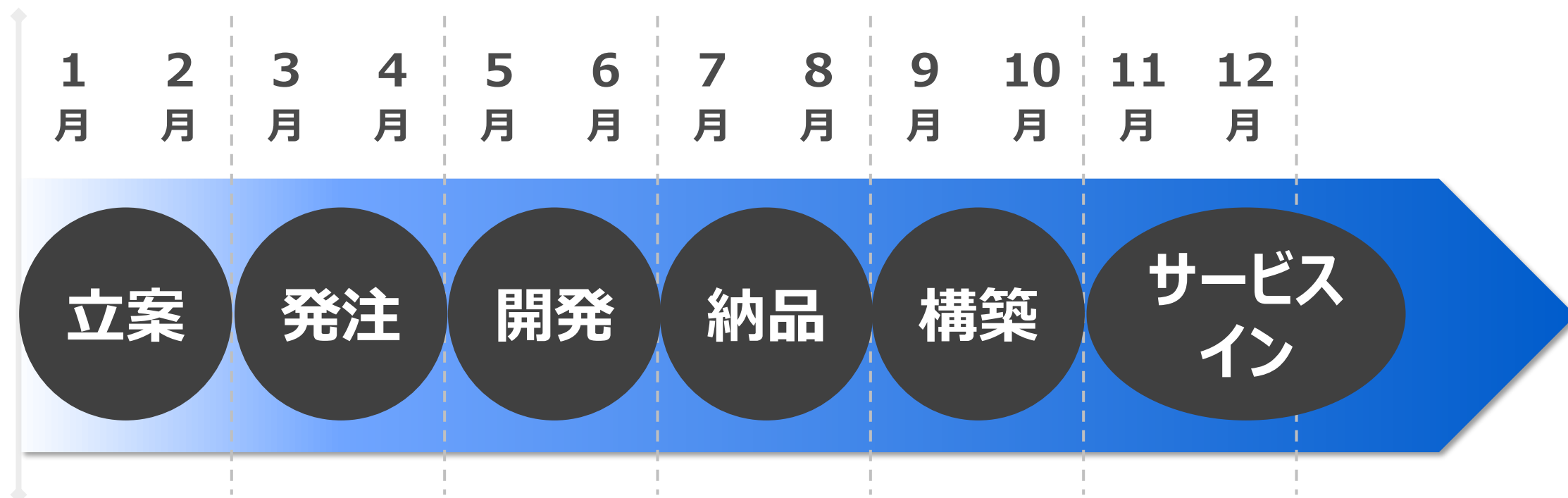
Q. インターコネクト
こんなに帯域いります？



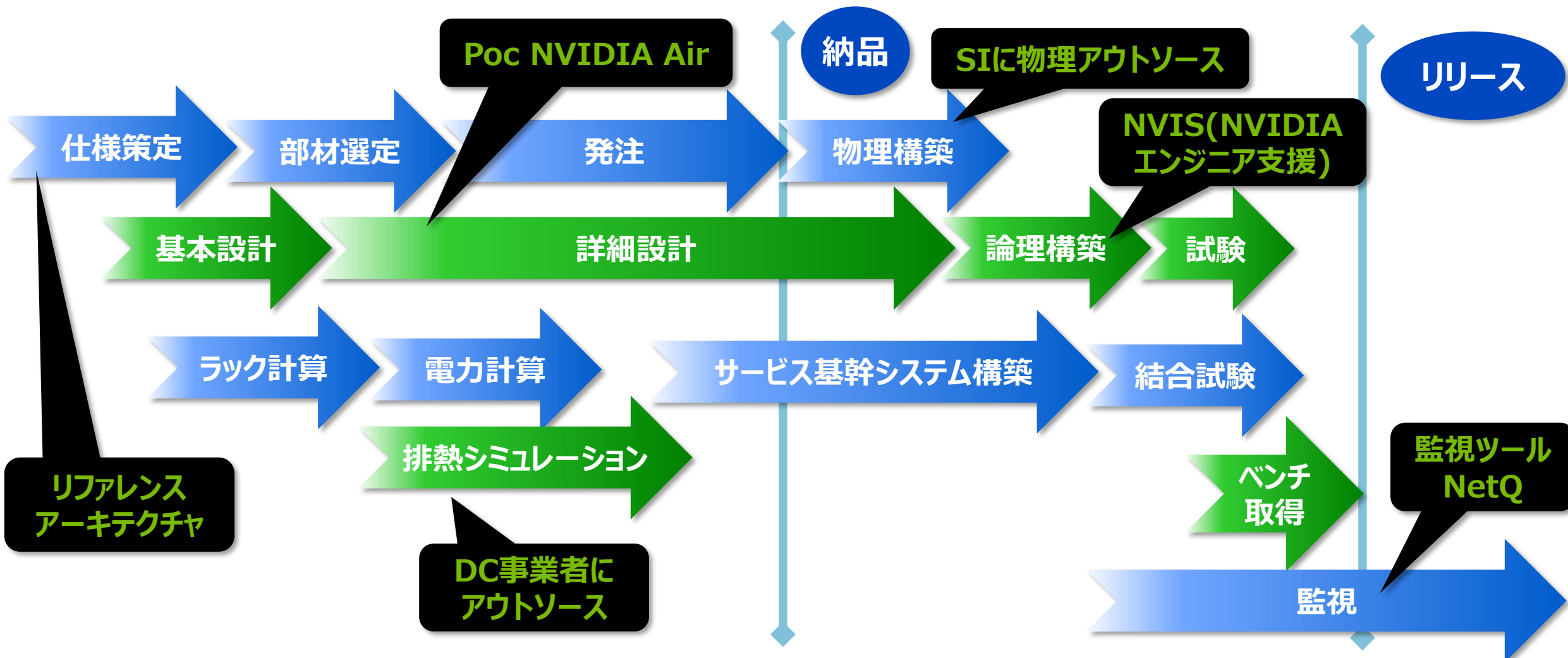
サービス提供は、
最短・最速

GPUは日に日に性能が上がるため **鮮度が命!!**

最短最速でサービスインする為の計画を策定



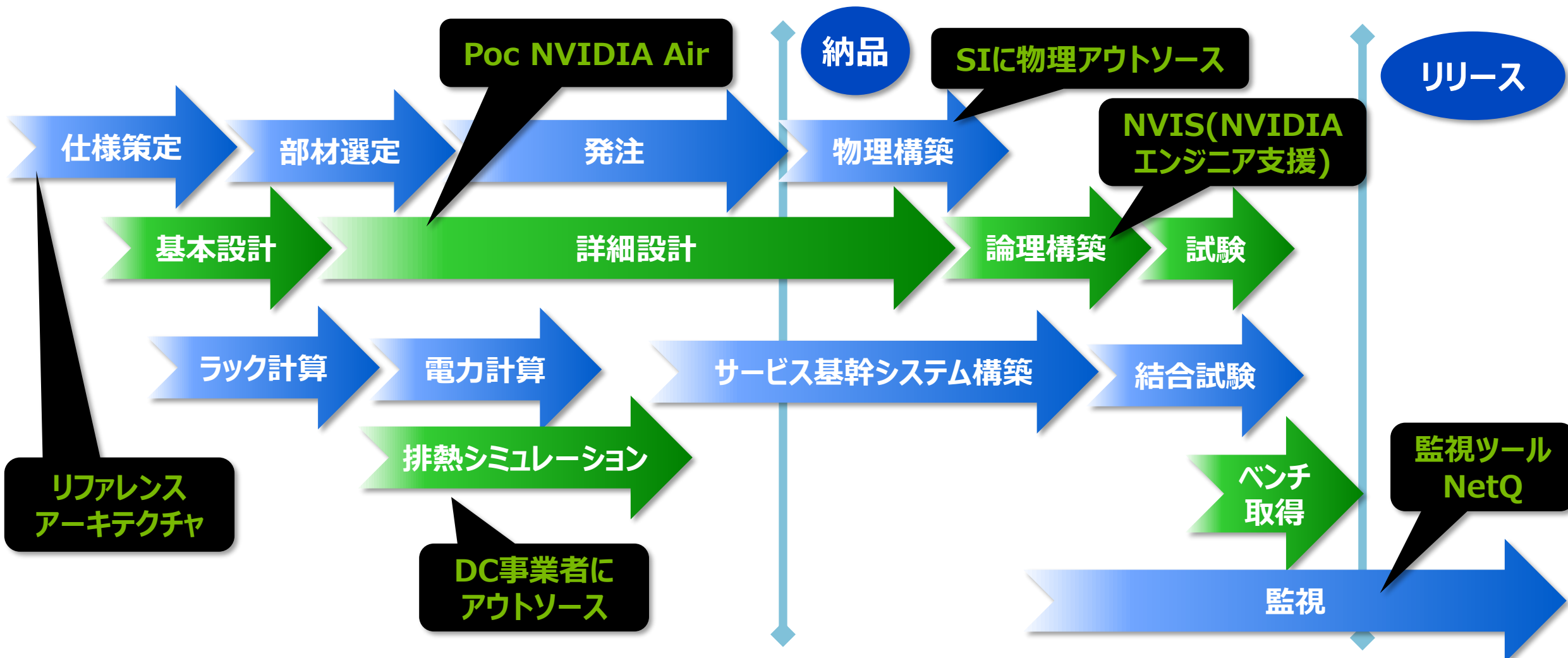
関係各社と連携、必要に応じてアウトソース、
NVIDIAフレームワーク活用。納品から3カ月でリリース！



GMOインターネットの事例

関係各社と連携、必要に
NVIDIAフレームワーク

Q. 短期リリースが大変！
上手なやり方ないの？



NVIDIA さんに聞いて みよう

1. 登壇者自己紹介
2. GMO GPUクラウドの構築事例紹介
 - 2-1. GMOインターネットの紹介
 - 2-2. システム構成
 - 2-3. ネットワーク構成
3. NVIDIAさんに聞いてみよう
 - 3-1. NVIDIAの紹介
 - 3-2. NVIDIAさんに質問をしてみよう
 - 3-3. 会場を交えての議論



主な事業分野

- GPU, AI, Network

2019年 Mellanox Technologies 買収

2020年 Cumulus Networks 買収

主なネットワーク製品

- Quantum InfiniBand プラットフォーム
- Spectrum Ethernet プラットフォーム
- ConnectX / BlueField NIC シリーズ
- LinkX トランシーバー&ケーブル ファミリー

InfiniBand を筆頭にスパコン、HPC 分野から続く
AI ネットワーキングにも欠かせないインターコネ
クト技術をリード

GPU ネットワーク = AI ネットワーク = インターコネクト

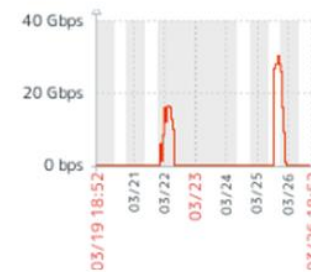
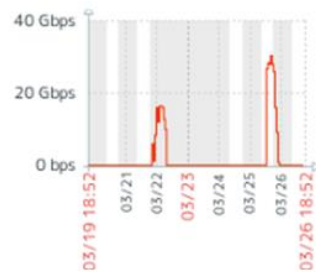
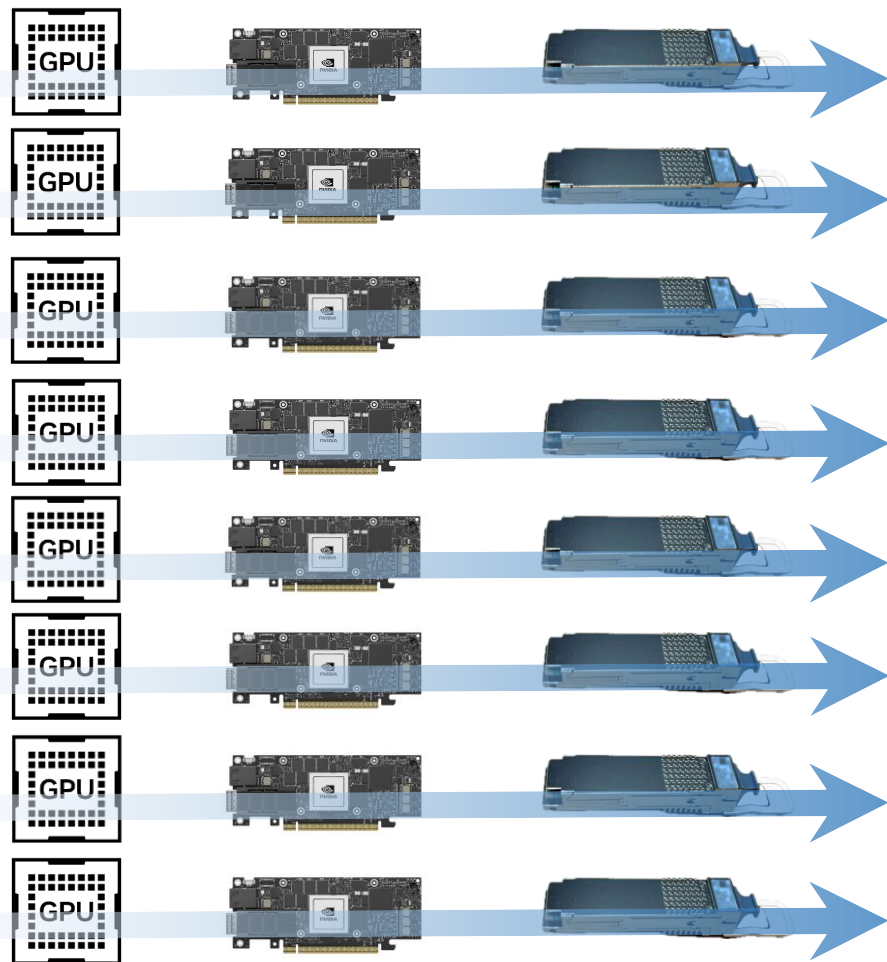
GPU = アクセラレーター

1Byte = 8bit

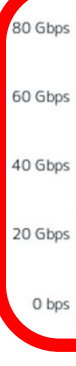
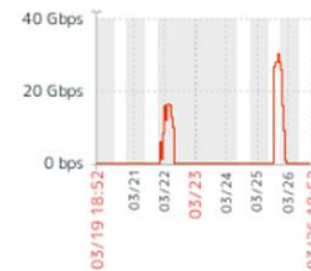
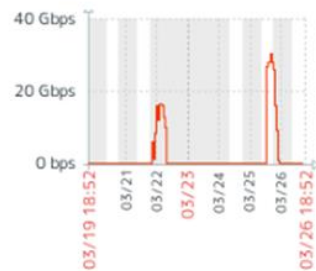
Q. インターコネクト
つてこんなに
帯域いります？

400Gラインで、実際の利帯域は多いときで、55Gほどに見える。

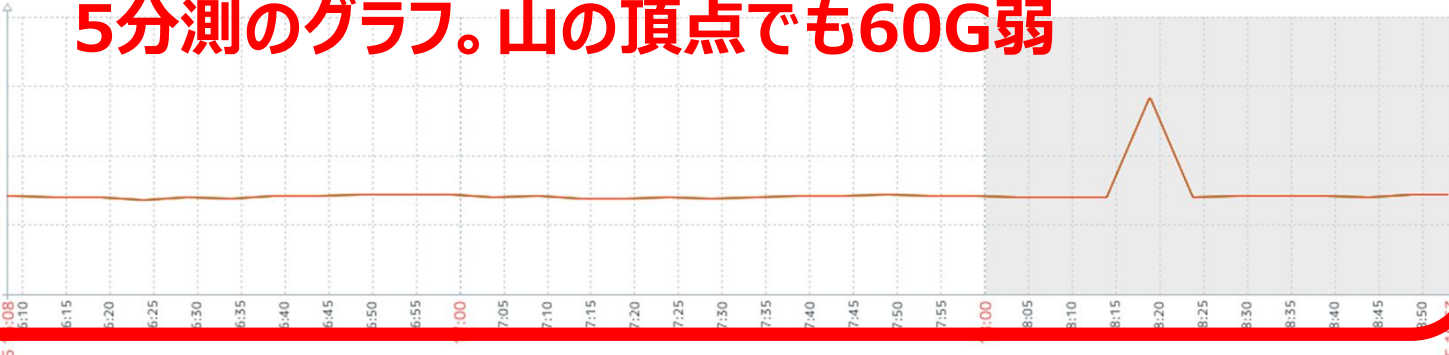
GPU BF3NIC 400Gb-Optics



一週間のグラフ



5分測のグラフ。山の頂点でも60G弱

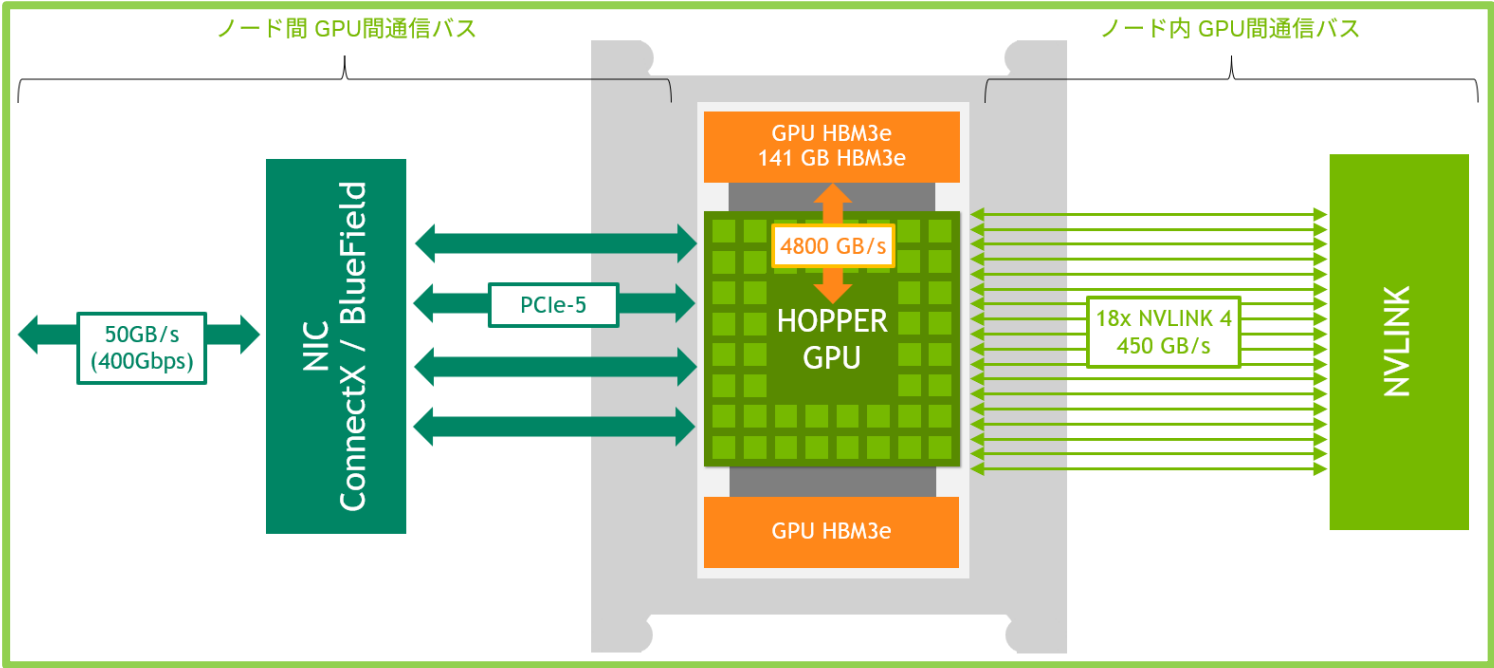


GMOの顧客だとフルスピード400Gは常に出てない

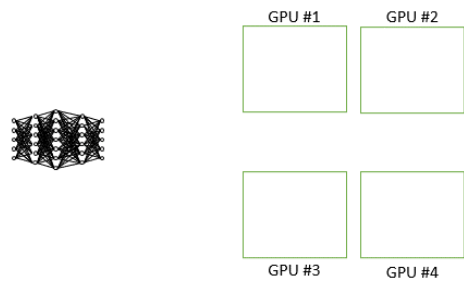
→実際のユーザで400G使うことはあるのでしょうか？

→低速だと何が困るのでしょうか？

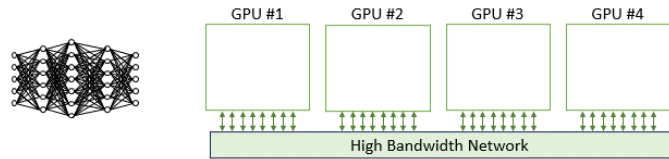
例えば100Gとか200Gだと、どうまずいのでしょうか？



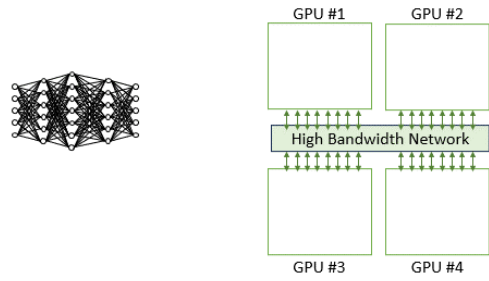
Data Parallel



Pipeline Parallel

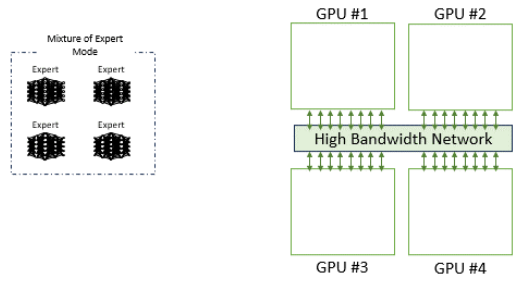


Tensor Parallel

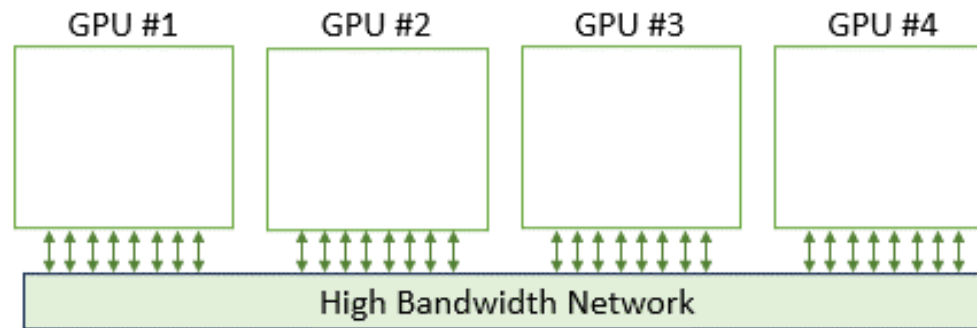
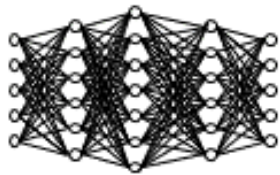


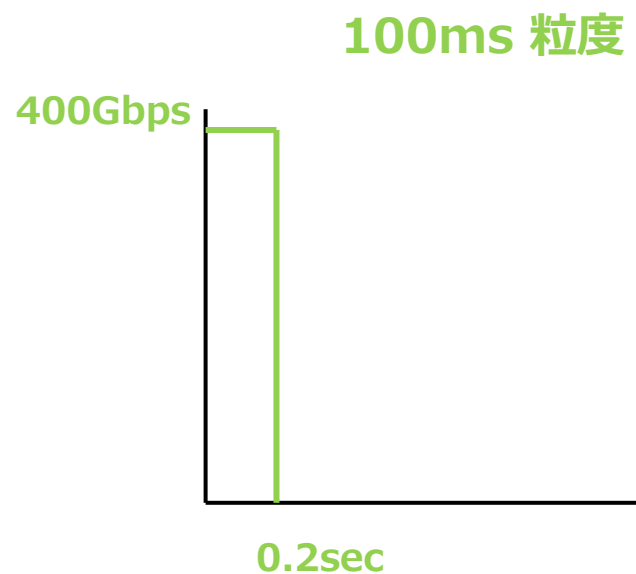
*C2C: Chip-to-Chip

Expert Parallel

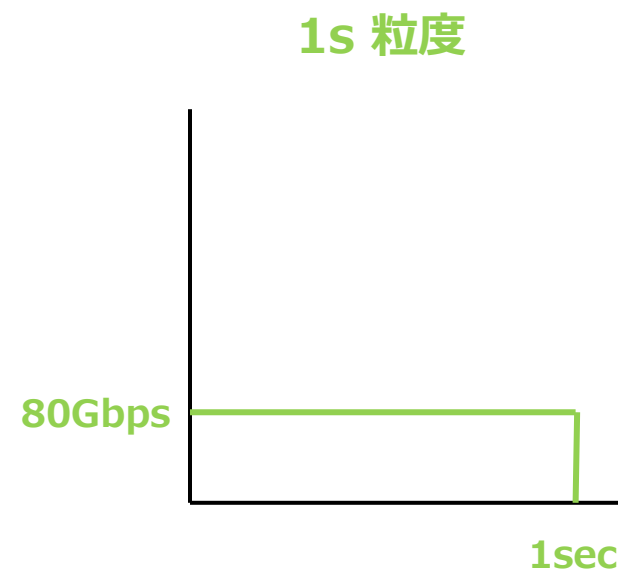


Pipeline Parallel





=



Q. **輻輳制御**って
そんなに大事なの？

RoCEv2 (RDMA over Converged Ethernet ver2) を利用し、RDMA通信のIBパケットをカプセル化

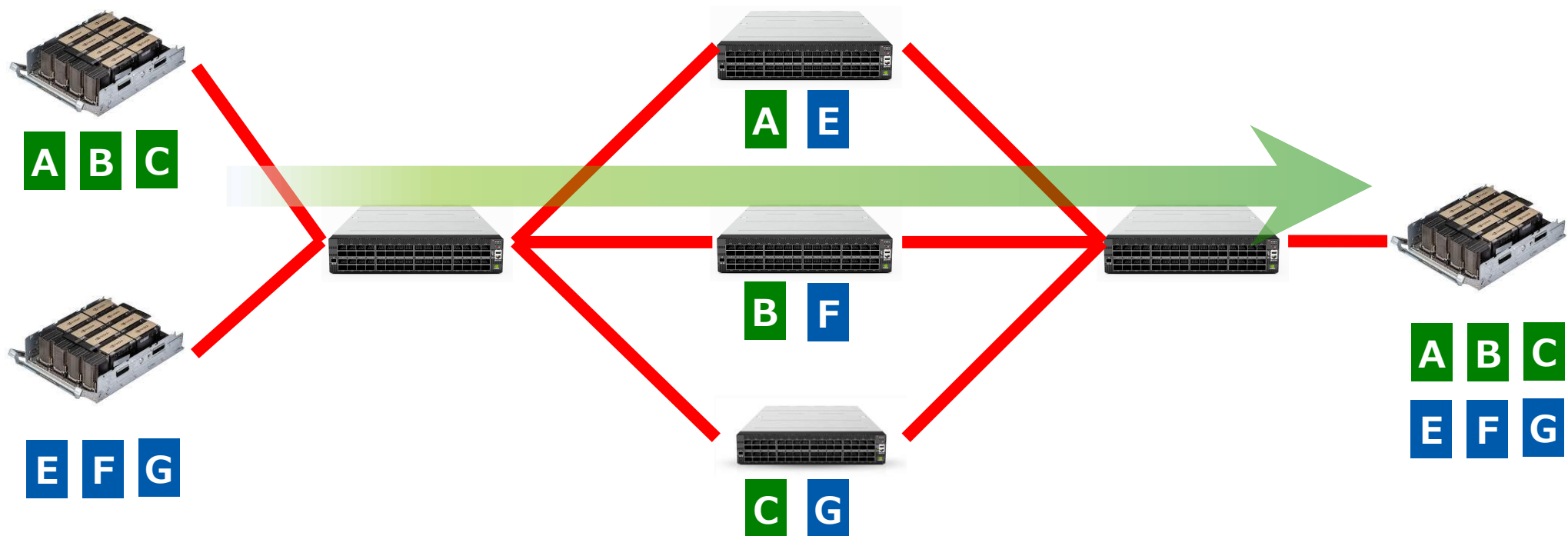
加えて、**ECN・PFC** などを利用し **輻輳制御**、
ロスレス通信を目指す。



※NVIDIA DOCS HUB RoCEv2 より引用

<https://docs.nvidia.com/networking/display/winofv55053000/rocev2>

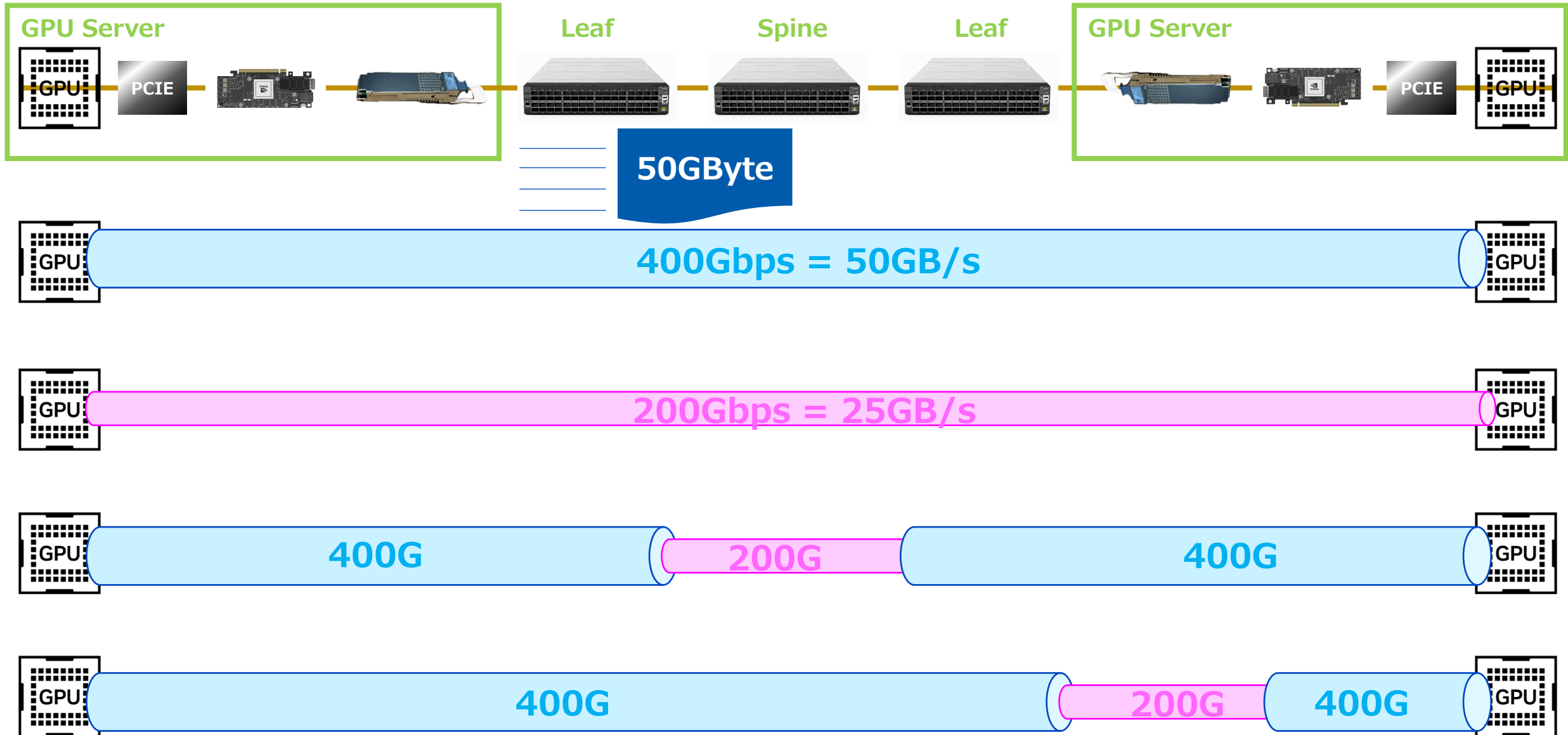
アダプティブブルーティングで全帯域を有効活用。
SWでインターフェイスごと通信を監視し、動的にบาลランシング
トラフィックを均等に転送し、
かつ、順序も維持することで**輻輳回避**を実現。



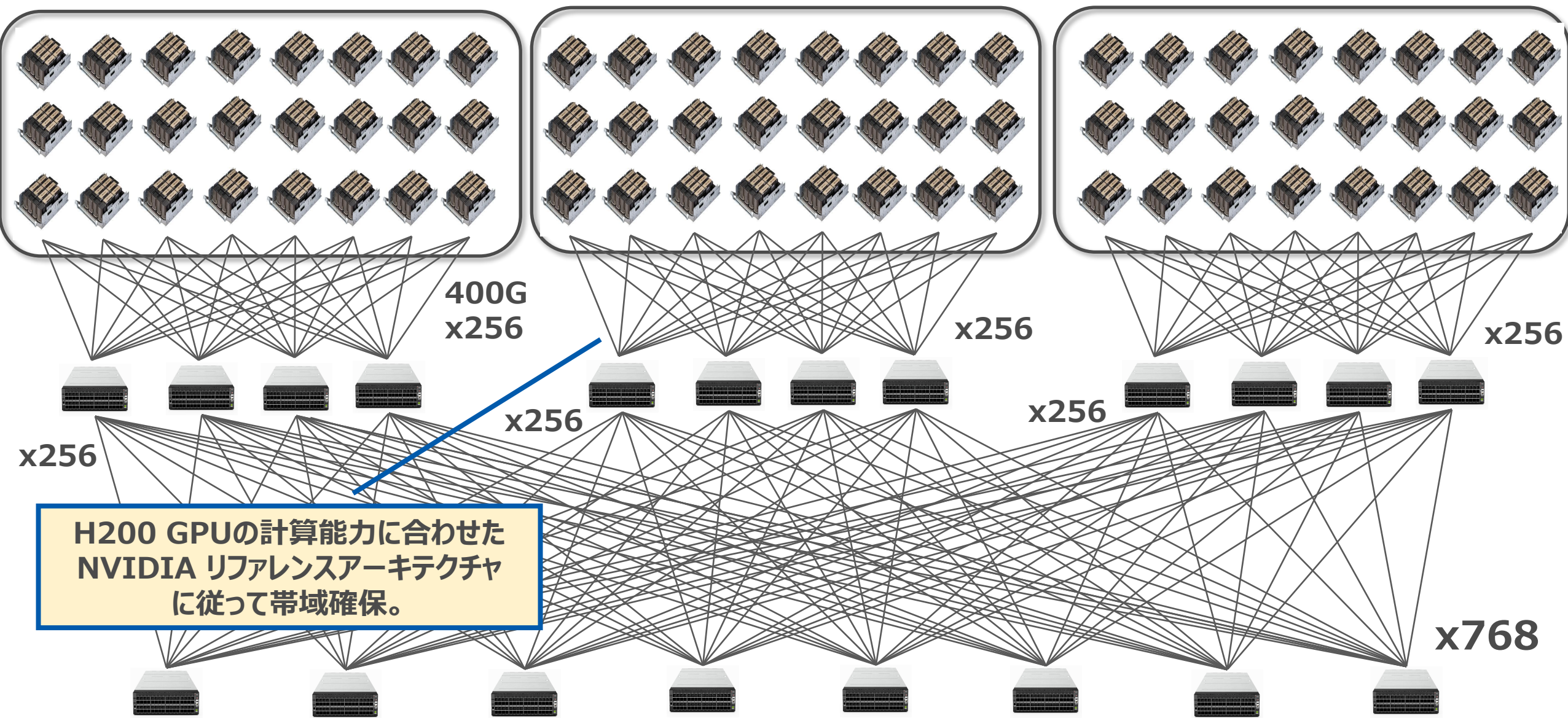
Spectrum-Switch & NVIDIA BlueField-3 SuperNIC を搭載した本システムでは、RoCEとアダプティブブルーディングに加えて**エンドツーエンドの輻輳制御**を実装。
GPU間通信の速度を調整し、より低遅延のロスレスNWを実現



- ・ RoCEとアダプティブブルーディングに加えて、
エンドツーエンドの輻輳制御を入れていますが、
 - 実装することで、どんな良い事があるのでしょうか？
 - 実装していないと、どういうシーンで困るのでしょうか？



Q. リファレンス
アーキテクチャ
の次の世代は？



今後、次世代のGPUを接続するインターコネクトネットワークの構築を検討しています

→次世代GPUのリファレンスアーキテクチャはどう変わりますか？

**→インターコネクトの帯域はどう変わりますか？
増えます？**

Q. Infinibandか
Ethernetか？

今後はどっちが主流？

弊社はマルチテナントなどを理由にEthernetを入れたけど、
Infinibandの場合のメリットも知りたい。

→正直なところ、
クラウド事業者にはどっちがおすすめですか？

■ 議論のポイント

- ・ 輻輳制御どうしていますか？
- ・ どのトポロジ採用してる？したい？
- ・ どれくらいの帯域を採用したい？
- ・ 構築の時短テクニックはありますか？
- ・ ICケーブルにDACとかAOC採用してる？
- ・ インターコネクト起因のパフォーマンス問題に直面したことはありますか？どのように切り分けましたか？

気になるところあれば、回答できる範囲で答えます！

すべての人にインターネット

GMO

