

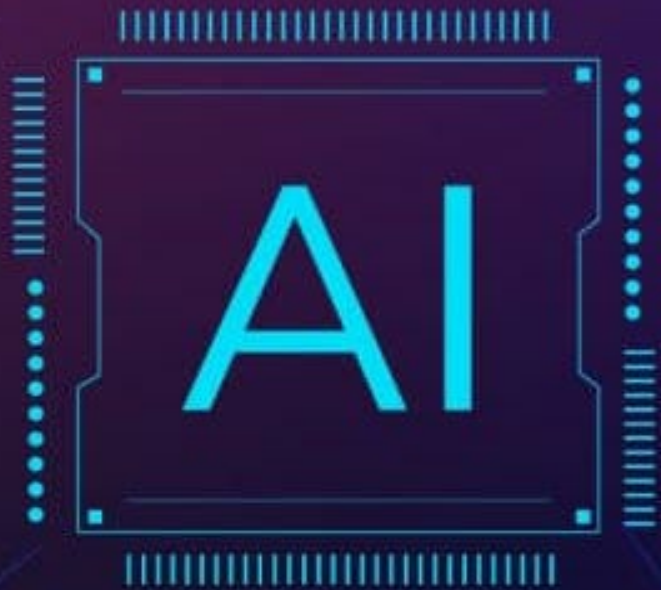
IA: Transformando la Arquitectura de los Data Centers

Antonio Cartagena

Gerente de Ventas NS, Centroamérica y el Caribe
Leviton



IA: Transformando la Arquitectura de Data Centers



Antonio Cartagena

Complex Challenges in the DC Space



Energía > Mayor potencia, mayor densidad



Cooling > Impacta una variedad de infraestructuras



Geografía > Nuevos territorios para el crecimiento



Latencia > Crítica para una operación eficiente de IA



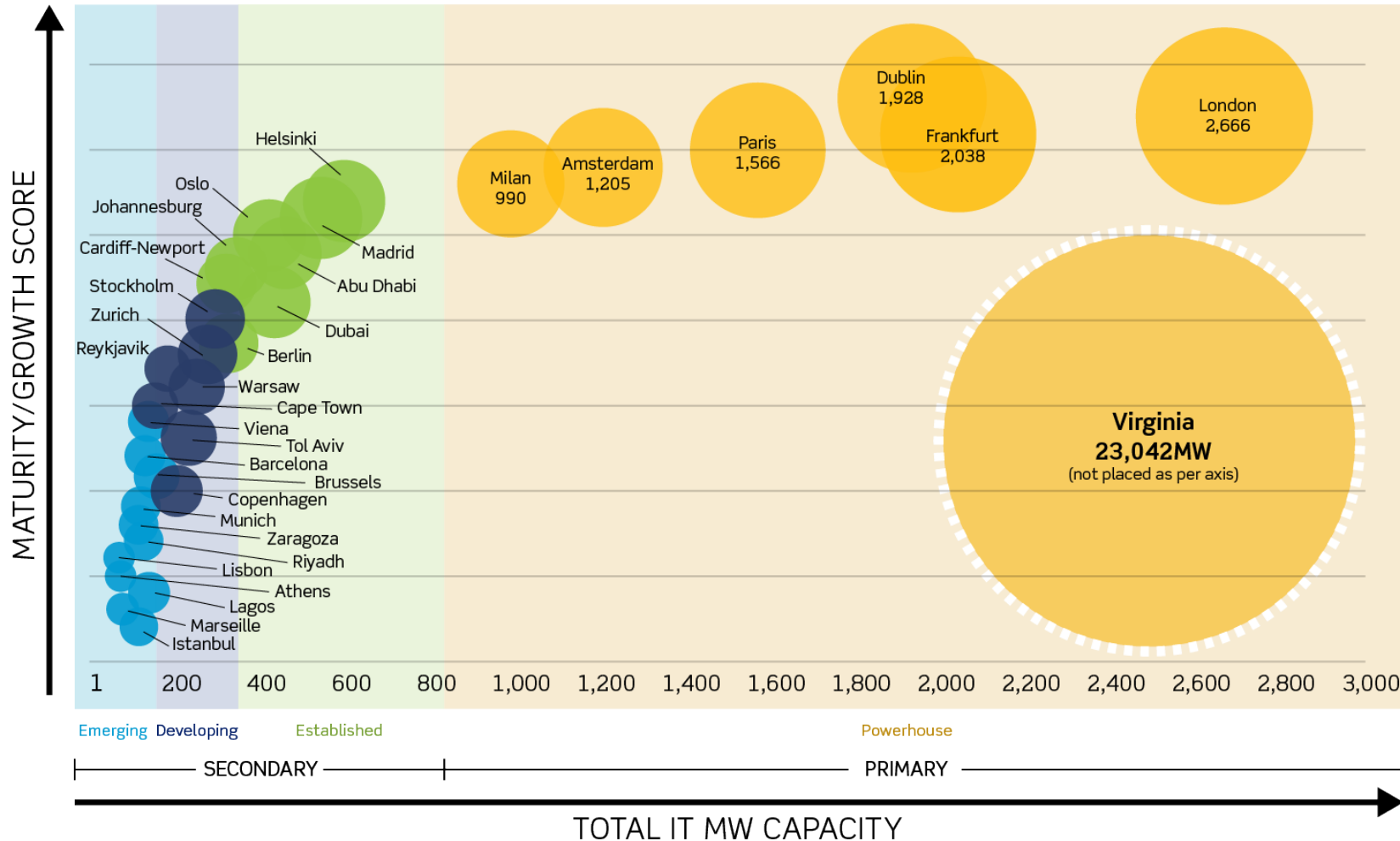
Velocidad de despliegue > ROI Rápido



Potencia Eléctrica

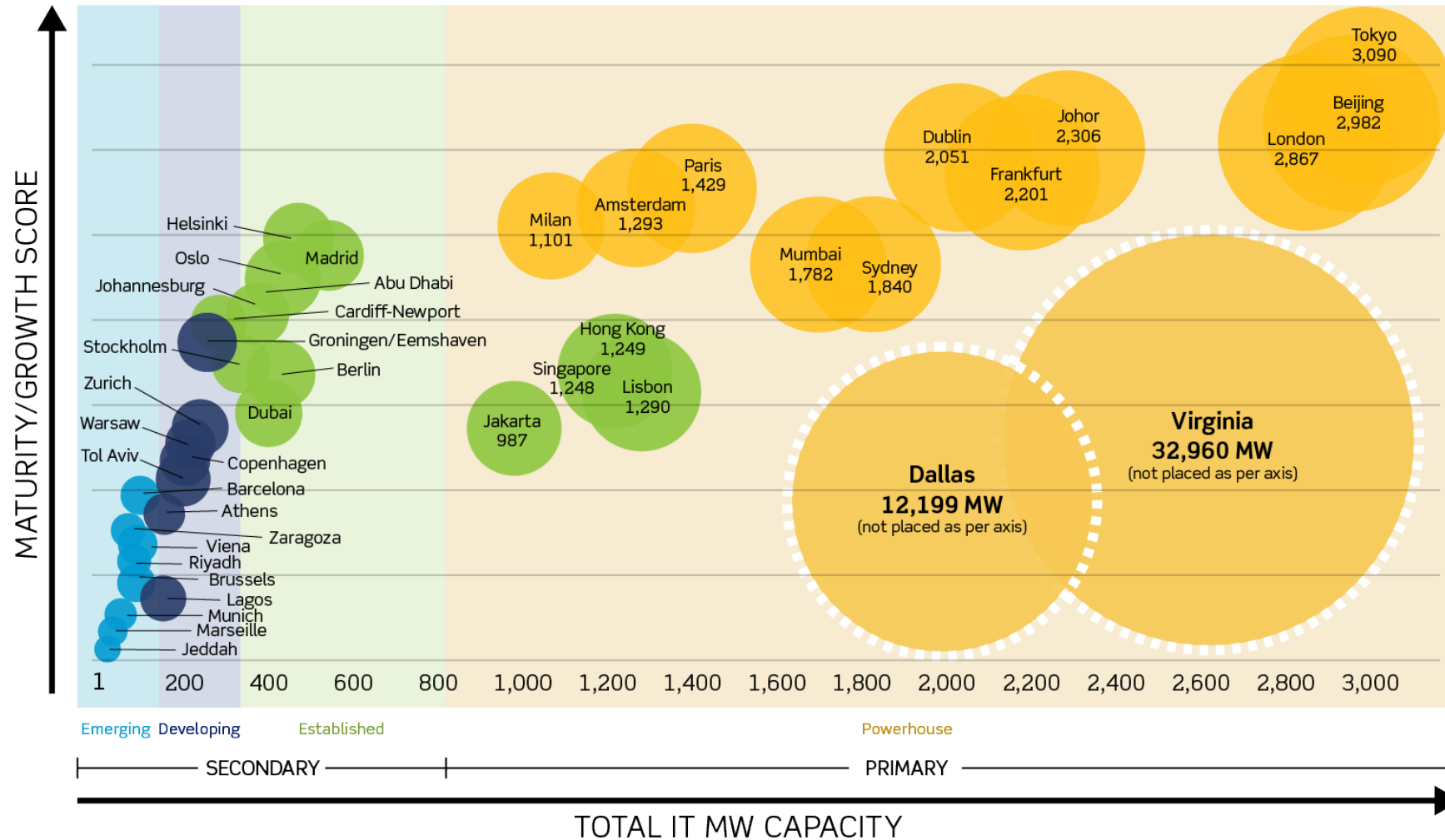
Nuevos territorios adquiriendo mayor importancia

2024



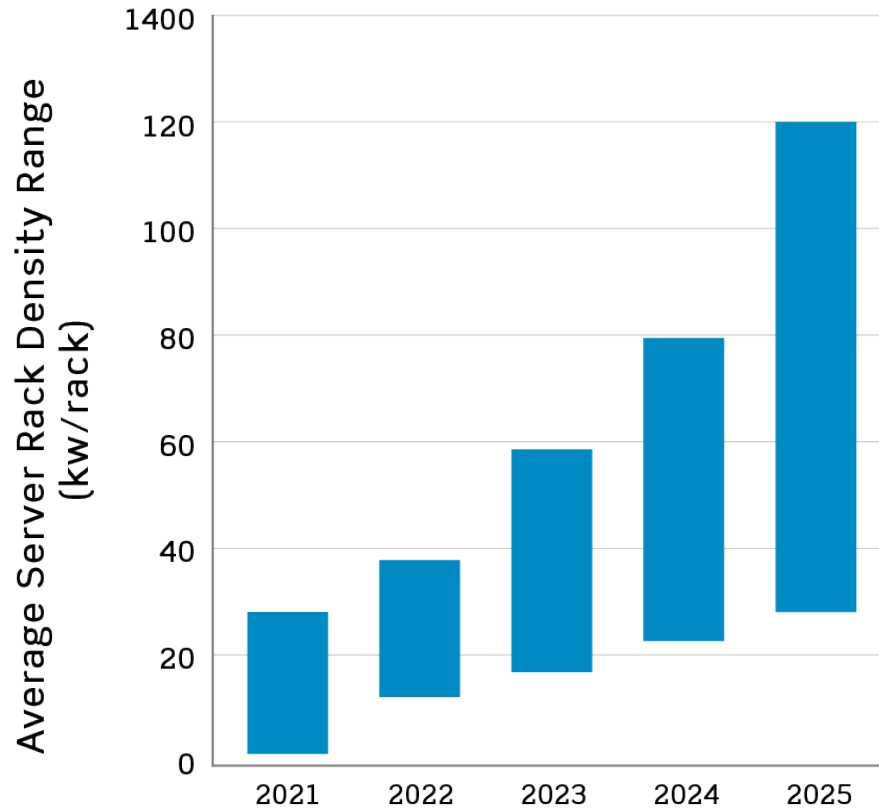
Nuevos territorios adquiriendo mayor importancia

2025

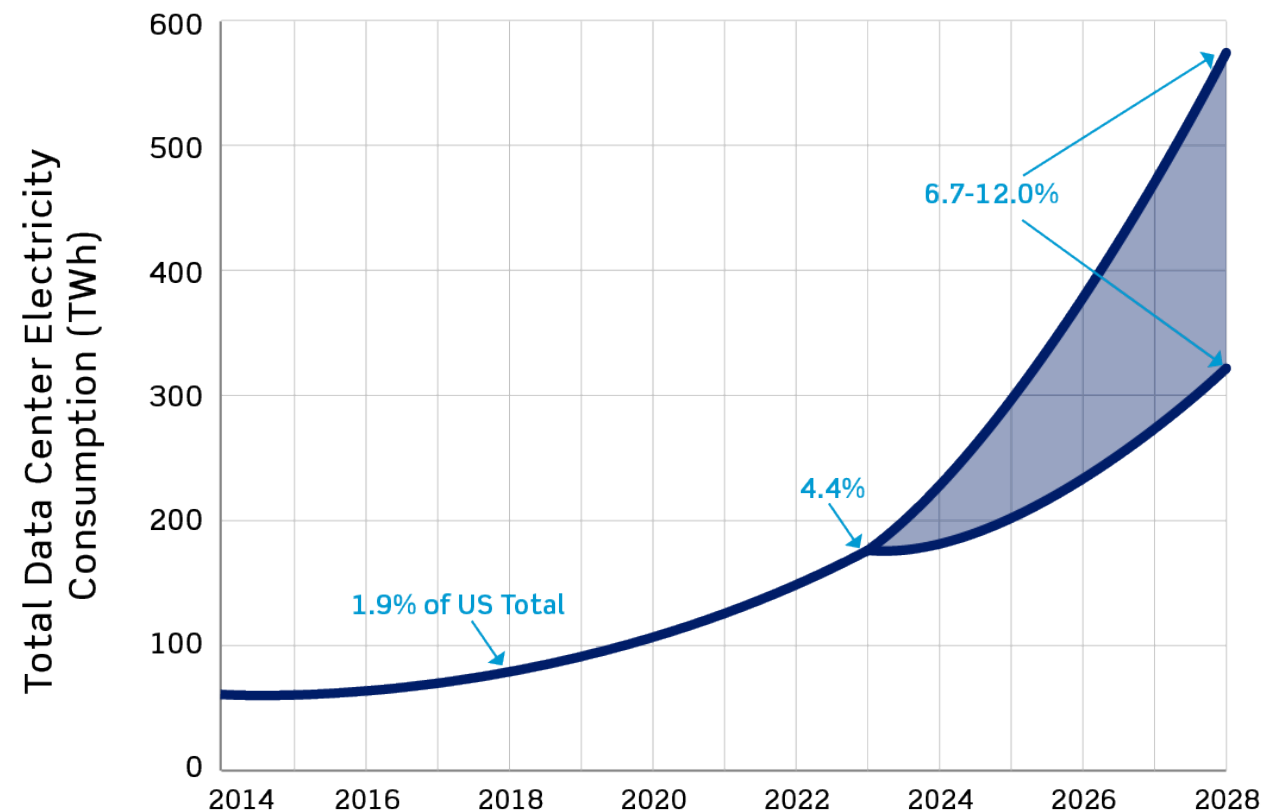


Consumo eléctrico de los Data Center

Crecimiento proyectado



Source: Cushman & Wakefield Research, Structure Research



Source: 2024 United States Data Center Energy Usage Report; Lawrence Berkeley National Lab

Emisiones de Gases de Efecto Invernadero

Data
Centre
3.7%



Transporte **3%**



Aerolíneas **2.5%**



Energía

¿Por qué GPU's?

La IA, el ML y el DL aprovechan las capacidades de procesamiento paralelo de una GPU. (Inteligencia Artificial, Aprendizaje Automático y Aprendizaje Profundo).

Miles de núcleos más en una GPU.

La CPU y la GPU trabajan juntas.

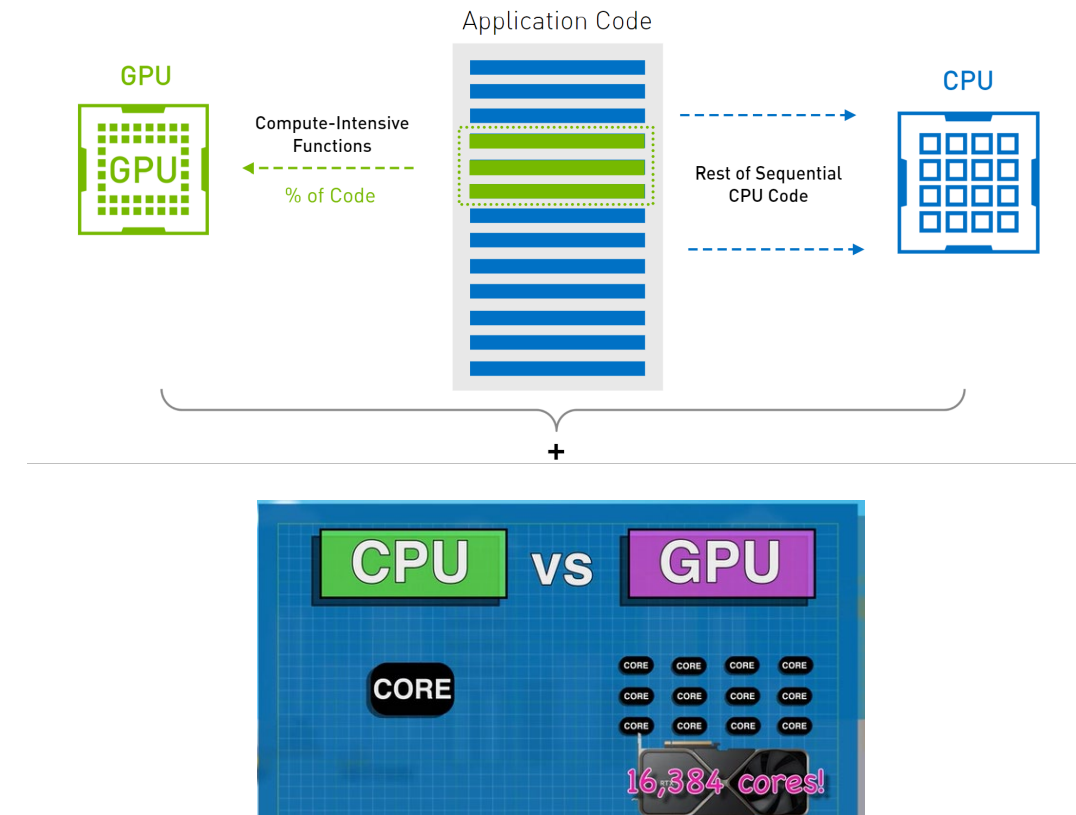
La CPU delega las tareas de cálculo intensivo a la GPU, de forma increíblemente rápida.

NVIDIA: Líder durante múltiples décadas en tecnología de GPU.

Videojuegos

Minería

How GPU Acceleration Works



NVIDIA no es el único fabricante de GPU's



Pero en 2024 NVIDIA tuvo el 90% del mercado de GPUs

NVIDIA DGX H100

“El estándar de oro para la infraestructura de IA”

Datasheet

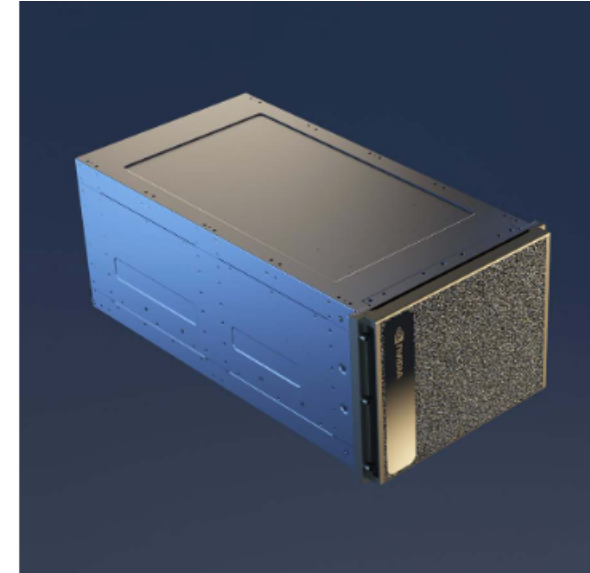


NVIDIA DGX H100

The gold standard for AI infrastructure.

Artificial intelligence has become the go-to approach for solving difficult business challenges. Whether improving customer service, optimizing supply chains, extracting business intelligence, or designing leading-edge products and services with generative AI and other transformer models, AI gives organizations across nearly every industry the mechanism to realize innovation. And as a pioneer in AI infrastructure, NVIDIA DGX™ provides the most powerful and complete AI platform for bringing these essential ideas to fruition.

NVIDIA DGX H100 powers business innovation and optimization. Part of the DGX platform and the latest iteration of NVIDIA's legendary DGX systems, DGX H100 is the AI powerhouse that's the



Specifications

GPU 8x NVIDIA H100 Tensor Core GPUs

GPU memory 640GB total

Performance 32 petaFLOPS FP8

NVIDIA® NVSwitch™ 4x

System power usage 10.2kW max

10.2 kW

CPU Dual Intel® Xeon® Platinum 8480C Processors
112 Cores total, 2.00 GHz (Base),
3.80 GHz (Max Boost)

Tecnología en Transceiver

Debido al corto alcance, los transceivers multimodo no siempre necesitan procesamiento de señal digital (DSP)

**Product Datasheet**

800G QSFP-DD SR8 Optical Transceiver
PN: VD-8CSR8CP-LP

Product Overview
VD-8CSR8CP-LP is a DSP free parallel 800G SR8 based 8-lane QSFP-DD pluggable transceiver that provides independent serial optical data links up to 8x 106 Gbps using PAM4 modulation format over multi-mode fiber.

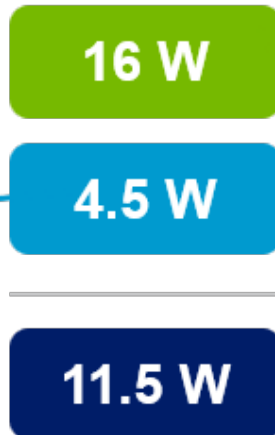
Features

- Up to 106 Gbps data rate per channel by PAM4 modulation
- Support 800GAUI-8 electrical interface
- Integrated 850nm VCSEL array and PD array without any DSP or CDR
- Single MPO16 connector receptacle optical interface compliant
- Hot-pluggable QSFP-DD form factor
- Low power consumption (4.5W)
- **Single +3.3V power supply**
- RoHS compliant
- IEEE 802.3cm, QSFP-DD HW Rev6.2, and CMIS 4.0 compliant

Applications

- Data Centers and Cloud Networks
- 800G Interconnect





EOLD-138HG-5H-M
Single-Mode, 800G, 8x100G QSFP-DD
With MPO-16 interface



Product Description
Eoptolink's QSFP-DD 8x100Gbps transceiver module can be used for 800 Gigabit Ethernet connections over 500m of single-mode fiber. The module includes eight parallel channels with a central wavelength of 1310nm, and the operating rate of each channel is 106.25Gbps. These 8-channel PAM4 parallel optical signals can be converted into 8-channel PAM4 electrical output signals; and there are 8 independent electrical input/output channels, which can convert PAM4 electrical input data into 8-channel PAM4 parallel optical signal. The transmitter of the module includes a bi-directional PAM4 re-timer ASIC and 8 EML Lasers. The receiver uses 8 photodiodes and two 4-channel TIA arrays, as well as the PAM4 re-timer.

Features

- Supports 850Gbps
- Single 3.3V Power Supply
- Up to 500m over SMF with KP4 FEC supported at the Host side
- MPO-16 connector
- 8x106.25Gbps (PAM4) electrical interface
- PIN and TIA array on the receiver side
- **Power dissipation < 16W**
- Case temperature range: 0°C to 70°C (commercial)

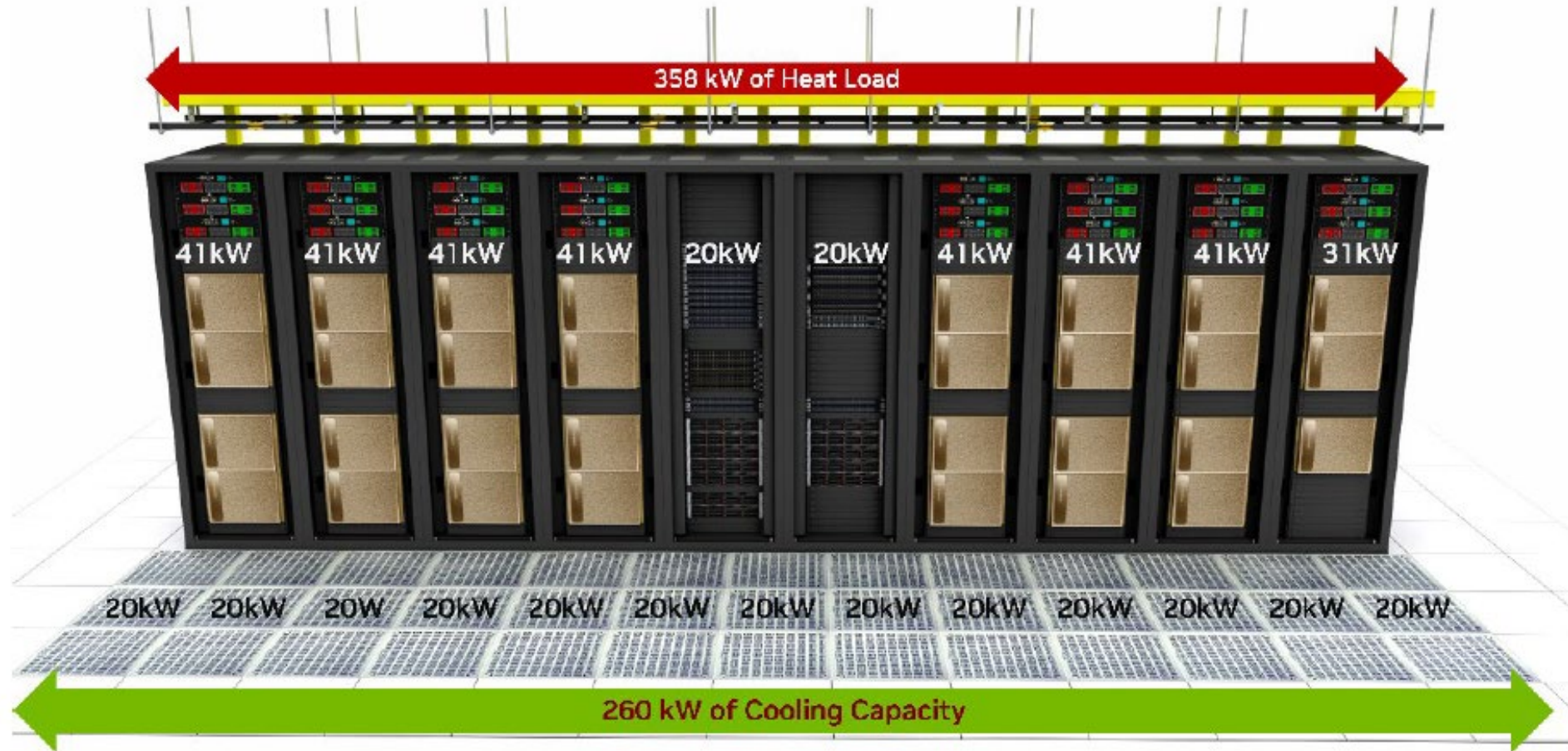
Applications*

- 8x100G Ethernet
- 2x400G Ethernet
- 1x800G Ethernet

Safety Certification: TUV/UL/FDA*
■ RoHS Compliant

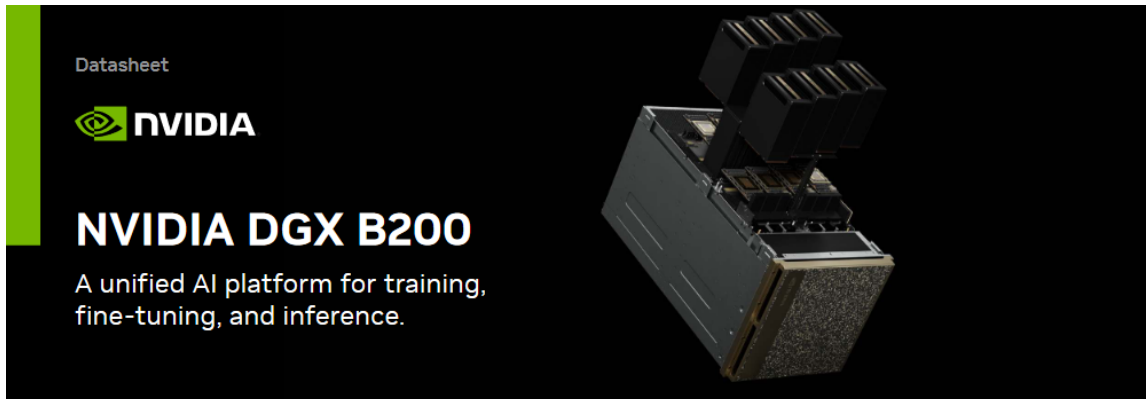
NVIDIA DGX Scalable Unit (SU)

Con switches Mgmt/Leaf in-row, refrigeración por aire



NVIDIA DGX B200

“The foundation for your AI center of excellence”



Powering the Next Generation of AI

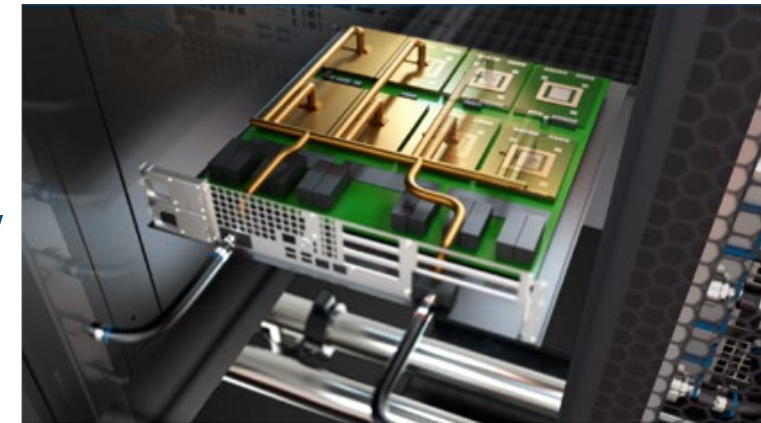
Artificial intelligence is transforming almost every business by automating tasks, enhancing customer service, generating insights, and enabling innovation. It's no longer a futuristic concept but a reality that's fundamentally reshaping how businesses operate. However, as AI workloads continue to develop, they're beginning to require significantly more compute capacity than most enterprises have available. To leverage AI, enterprises need high-performance computing, storage, and networking capabilities that are secure, reliable, and efficient.

Enter **NVIDIA DGX™ B200**, the latest addition to the **NVIDIA DGX platform**. This unified AI platform defines the next chapter of generative AI by taking full advantage of NVIDIA Blackwell GPUs and high-speed interconnects. Configured with eight Blackwell GPUs, DGX B200 delivers unparalleled generative AI performance with a massive 1.4 terabytes (TB) of GPU memory and 64 terabytes per second (TB/s) of HBM3e memory bandwidth, and 14.4 TB/s of all-to-all GPU bandwidth, making it uniquely suited to handle any enterprise AI workload.

Key Features

NVIDIA DGX B200

- > Built with eight NVIDIA Blackwell GPUs
- > 1.4TB of GPU memory space
- > 72 petaFLOPS of training performance
- > 144 petaFLOPS of inference performance
- > NVIDIA networking
- > Dual 5th generation Intel® Xeon® Scalable Processors
- > Foundation of NVIDIA DGX



Direct Liquid Cooling

DGX B200 Technical Specifications	
GPU	8x NVIDIA Blackwell GPUs
GPU Memory	1,440GB total, 64TB/s HBM3e bandwidth
Performance	72 petaFLOPS FP8 training and 144 petaFLOPS FP4 inference
NVIDIA® NVSwitch™	2x
NVIDIA NVLink Bandwidth	14.4 TB/s aggregate bandwidth
System Power Usage	~14.3kW max
CPU	2 Intel® Xeon® Platinum 8570 Processors 112 Cores total, 2.1 GHz (Base), 4 GHz (Max Boost)
System Memory	2TB, configurable to 4TB
Networking	4x QSFP ports serving 8x single-port NVIDIA ConnectX-7 VPI > Up to 400Gb/s InfiniBand/Ethernet 2x dual-port QSFP112 NVIDIA BlueField-3 DPU > Up to 400Gb/s InfiniBand/Ethernet

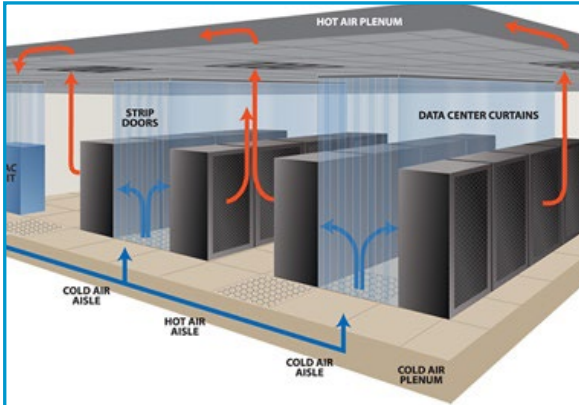
14.3 kW

Source: Nvidia



Enfriamiento

Tecnologías de refrigeración en los Data Center



Aire a aire

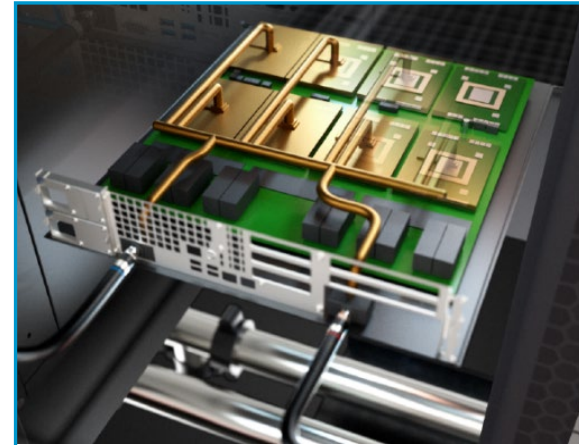
~1's of kW

Examples: REHVA, ASHRAE



Intercambiador de calor
de puerta trasera

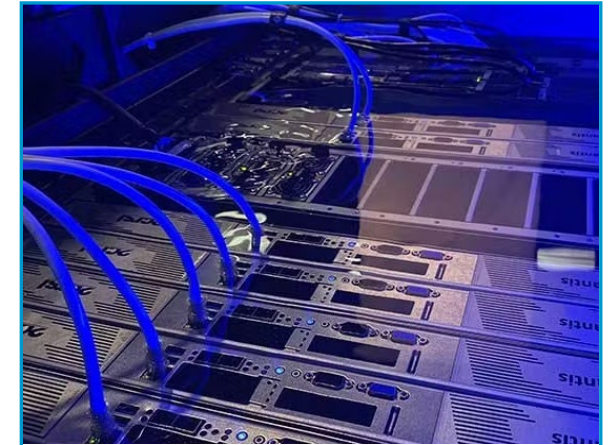
~10's of kW



Refrigeración líquida directa

Up to ~100's of kW

Vendor Specific



Refrigeración por inmersión

> 100's of kW

OCP

Consumo de energía de la IA

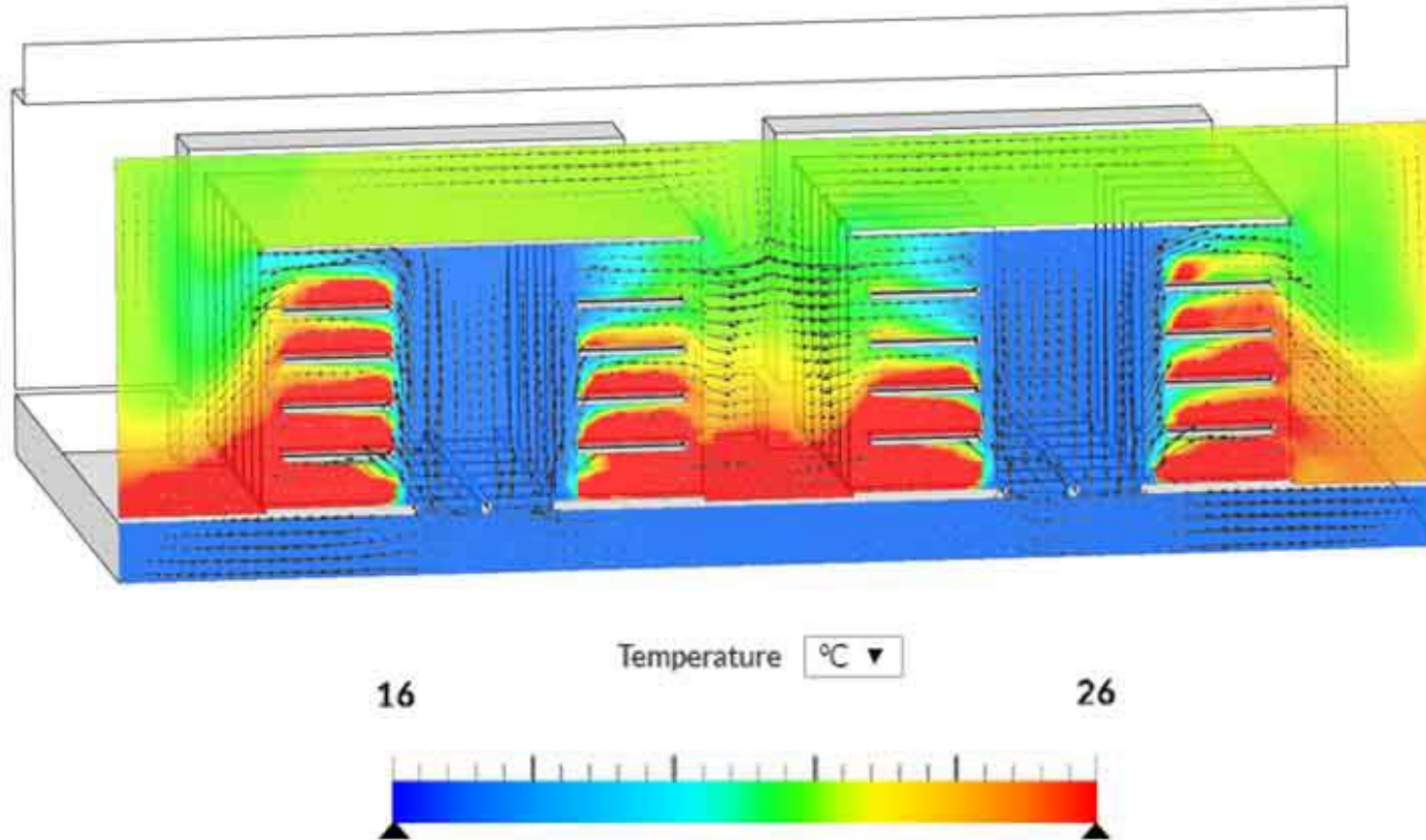
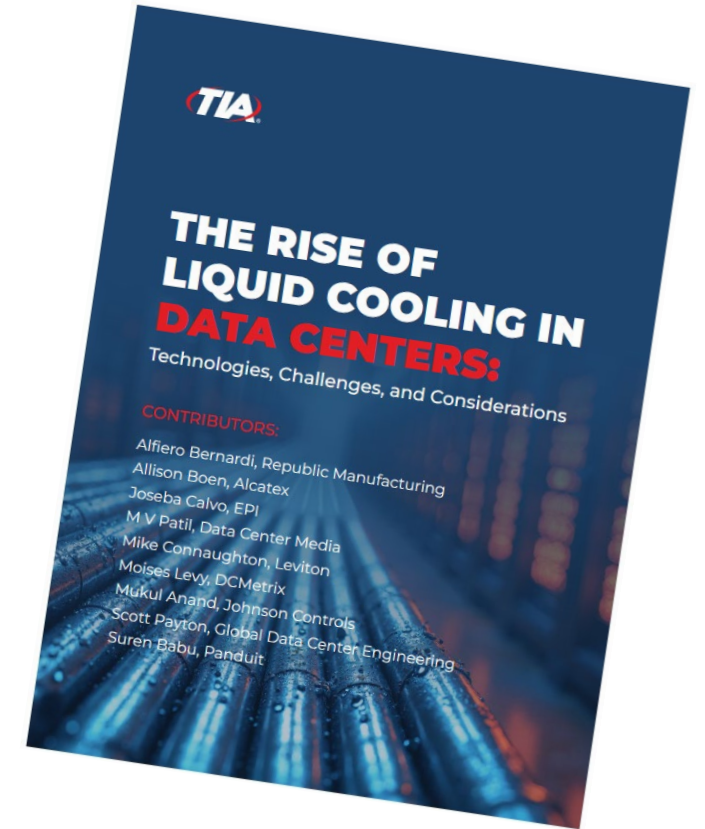


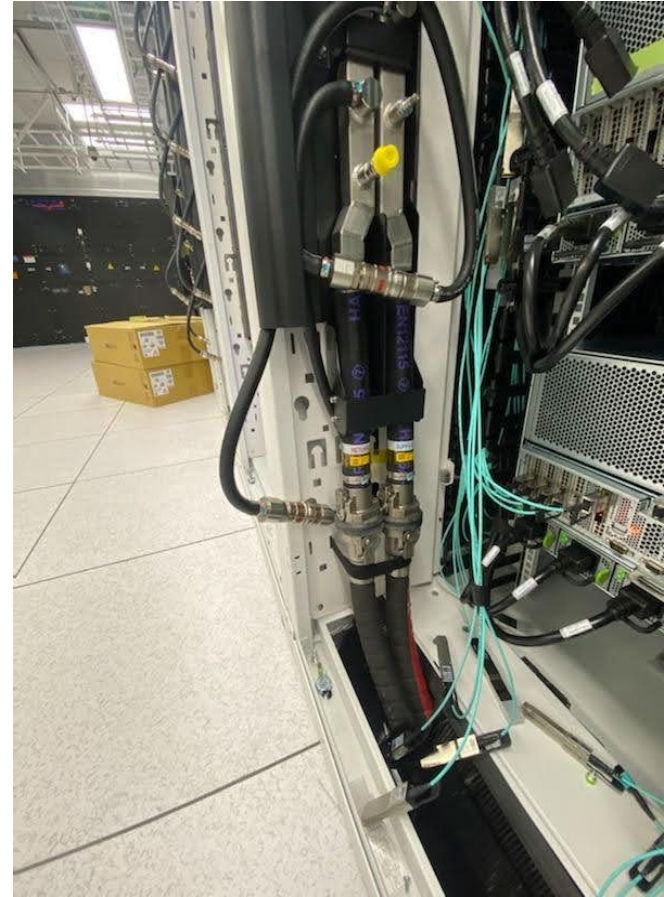
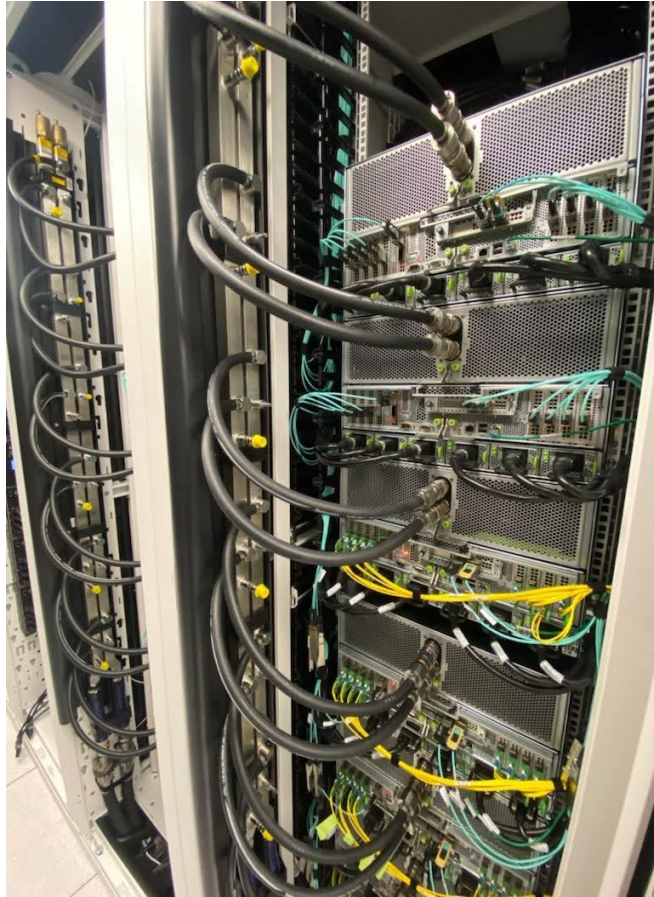
Image: Sim Scale

Refrigeración líquida

- TIA White Paper
 - El auge de la refrigeración líquida en los Data Centers
 - Tecnologías, desafíos y consideraciones
- Consideraciones sobre carga en el suelo
 - Aumento significativo para immersion cooling
- Arquitectura
 - Las acometidas que tradicionalmente se utilizaban para cableado ahora también para tuberías.



Direct to Chip, ya se está haciendo



Refrigeración líquida

Liquid Assisted Cooling

- Mínimo impacto en el cableado

Direct to Chip Cooling

- Mínimo impacto en el cableado

Fluid Immersion

- La compatibilidad del cable es una preocupación



Green Revolution Immersion Cooling

Cableado certificado para inmersión de medios líquidos

- Las soluciones de cobre y fibra están preparadas para la inmersión en fluidos de refrigeración seleccionados, habiendo sido probadas y validadas según los requisitos de compatibilidad de materiales en refrigeración por inmersión de la OCP (Open Compute Project), Rev 1.0.
- líder en el área de pruebas de compatibilidad de fluidos. Ningún otro competidor de cableado estructurado ha realizado un anuncio de este tipo.



Redes de Inteligencia Artificial y Machine Learning



Hyperscale y IA Data Centers

Todos los Data Centers
Hyperscale tienen

AI Deployments

Pero **no todos** los

AI Deployments

se consideran Hyperscale



Las velocidades en las redes de IA siguen evolucionando

Publicado en 2022

- IEEE 802.3ck especifica 100 Gb/s, 200 Gb/s, y 400 Gb/s
 - Interfaces eléctricas basadas en la señalización de 100 Gb/s

Publicado en 2024

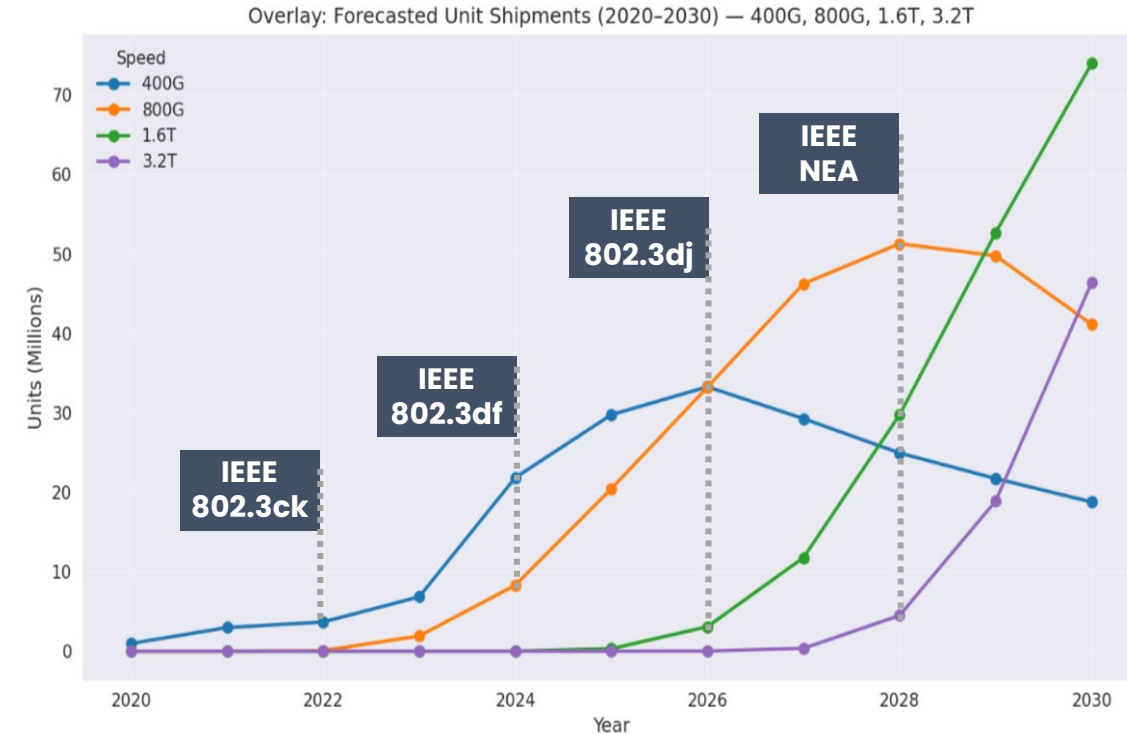
- IEEE 802.3df especifica 400 y 800 Gb/s con señalización de 100 Gb/s

Próximo a finalizar

- IEEE 802.3dj especifica 200–1600 Gb/s con señalización de 200Gb/s
 - Actualmente en la fase de borrador 3

Próximamente

- La próxima aplicación Ethernet de IEEE especificará la señalización de 400 Gb/s para 3,2 Tb/s.
 - Inicialmente, está diseñada para ópticas integradas.

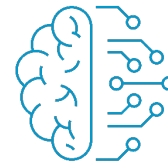
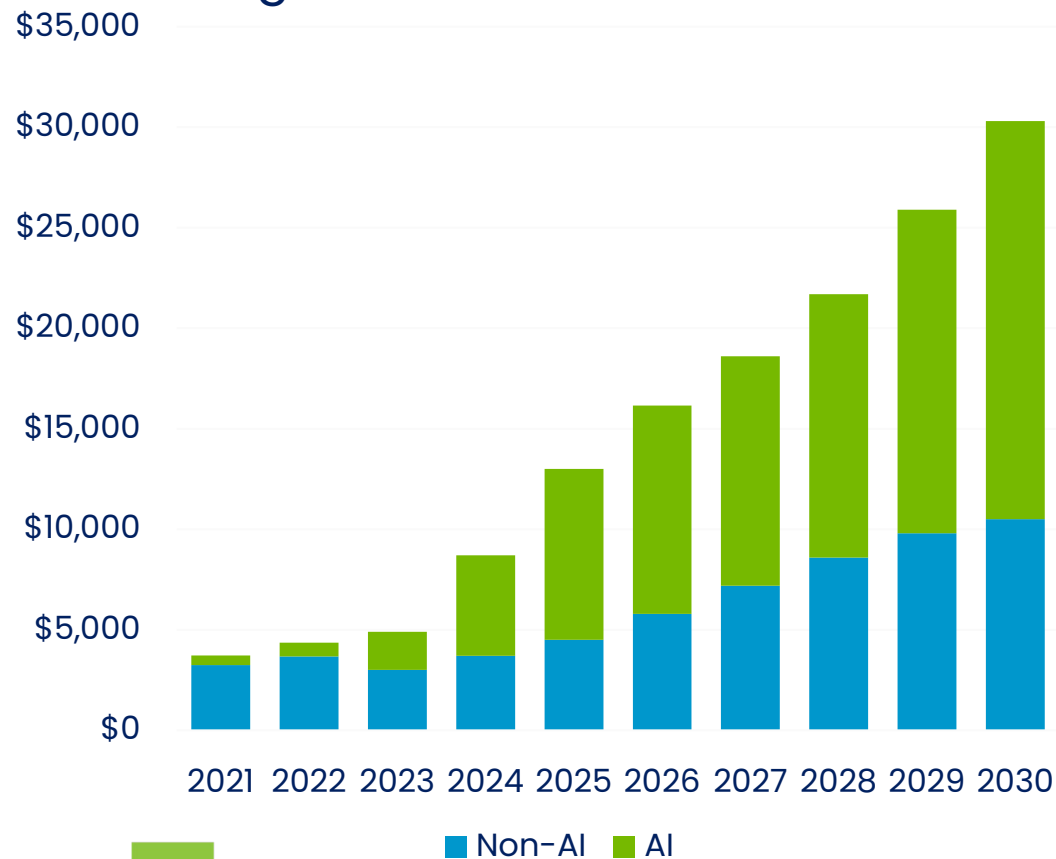


Source: LightCounting

Crecimiento de transceivers

El mayor motor de crecimiento actual

Significant Driver Port Growth

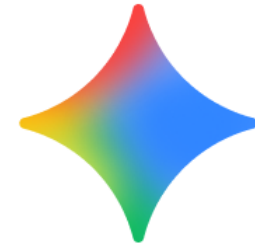


Crecimiento de la IA y Machine Learning en los Data Center

% of Transceiver Sales			
	2020		2028 Projection in 2025
AI Clusters	15%	➔	58%
Resto de Redes Cloud	85%	➔	42%



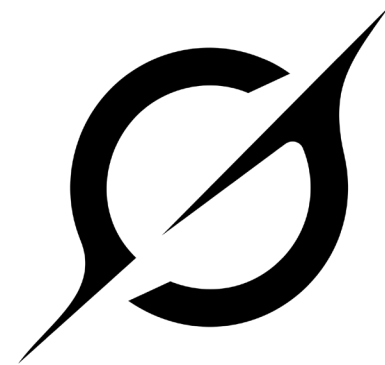
ChatGPT



Gemini



Copilot



Grok



Claude



Meta AI



Certified
Associate

**AI Infrastructure
and Operations**

➤ La multinacional tecnológica **NVIDIA** utiliza la infraestructura de fibra óptica y cobre de **Leviton** en sus oficinas y centros de datos de todo el mundo.

20,000

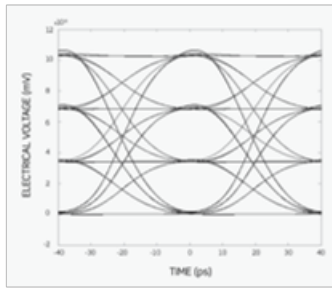
Productos disponibles para
NVIDIA en EE.UU., Reino Unido,
Dubái y Hong Kong

Cuestiones que deben de abordarse

Tres caminos hacia velocidades de datos más altas



Velocidad del carril
100G, 200G, 400G, ...



1

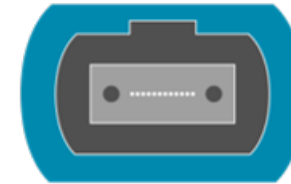
Número de longitudes de onda
1, 4, 8, ...

2

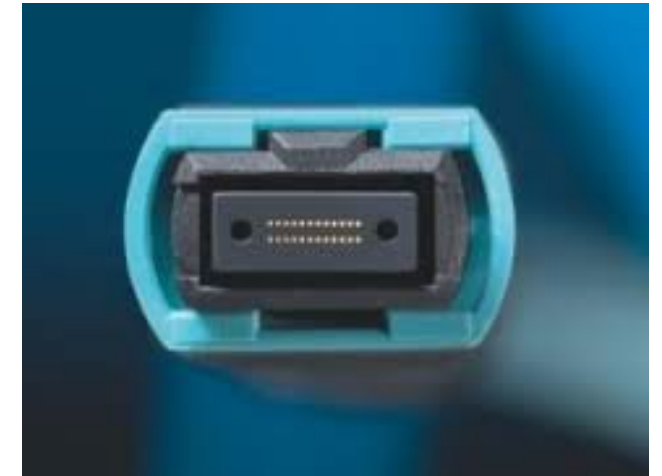
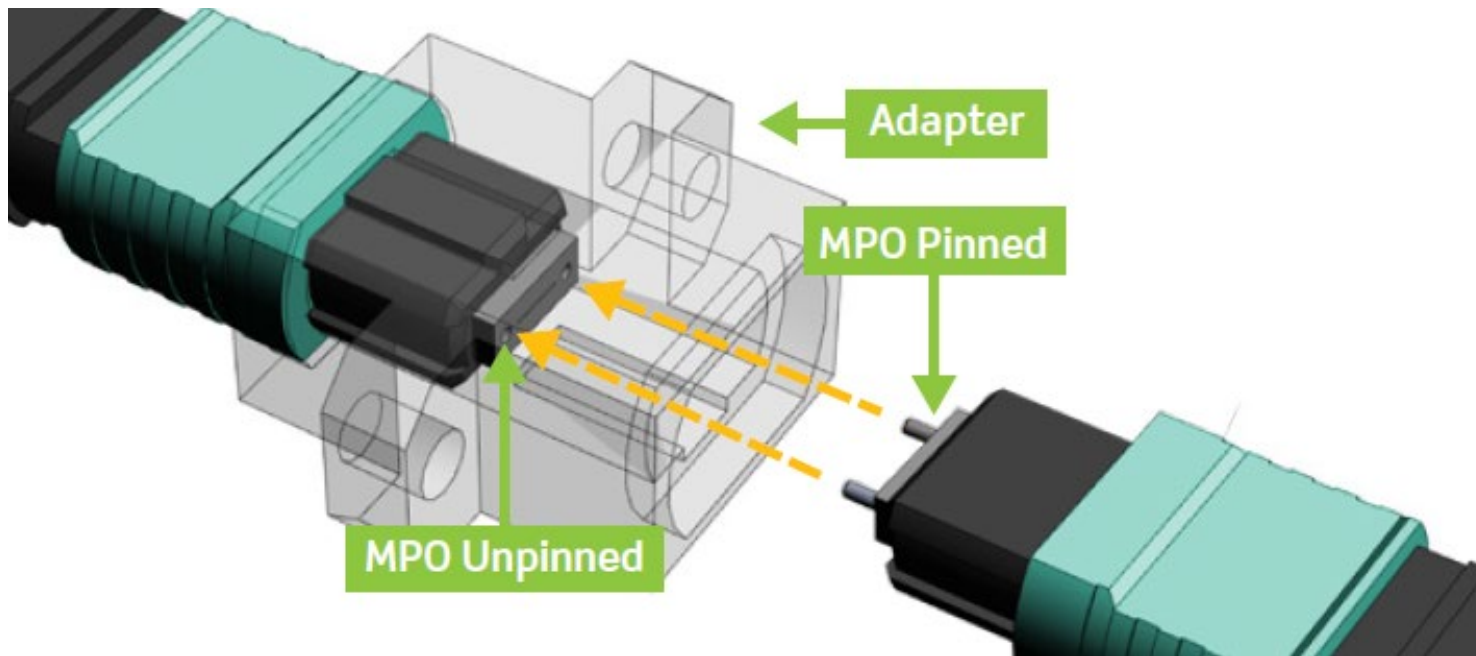


3

Recuento de fibras
2, 8, 16, ...



Conector MTP y/o MPO



Transmisión Multilínea

EEE 802.3ba (Junio 2010)

- **Principios Básicos**

- Transmisión multilínea
- conector MPO (IEC 61754-7)

- **40G**

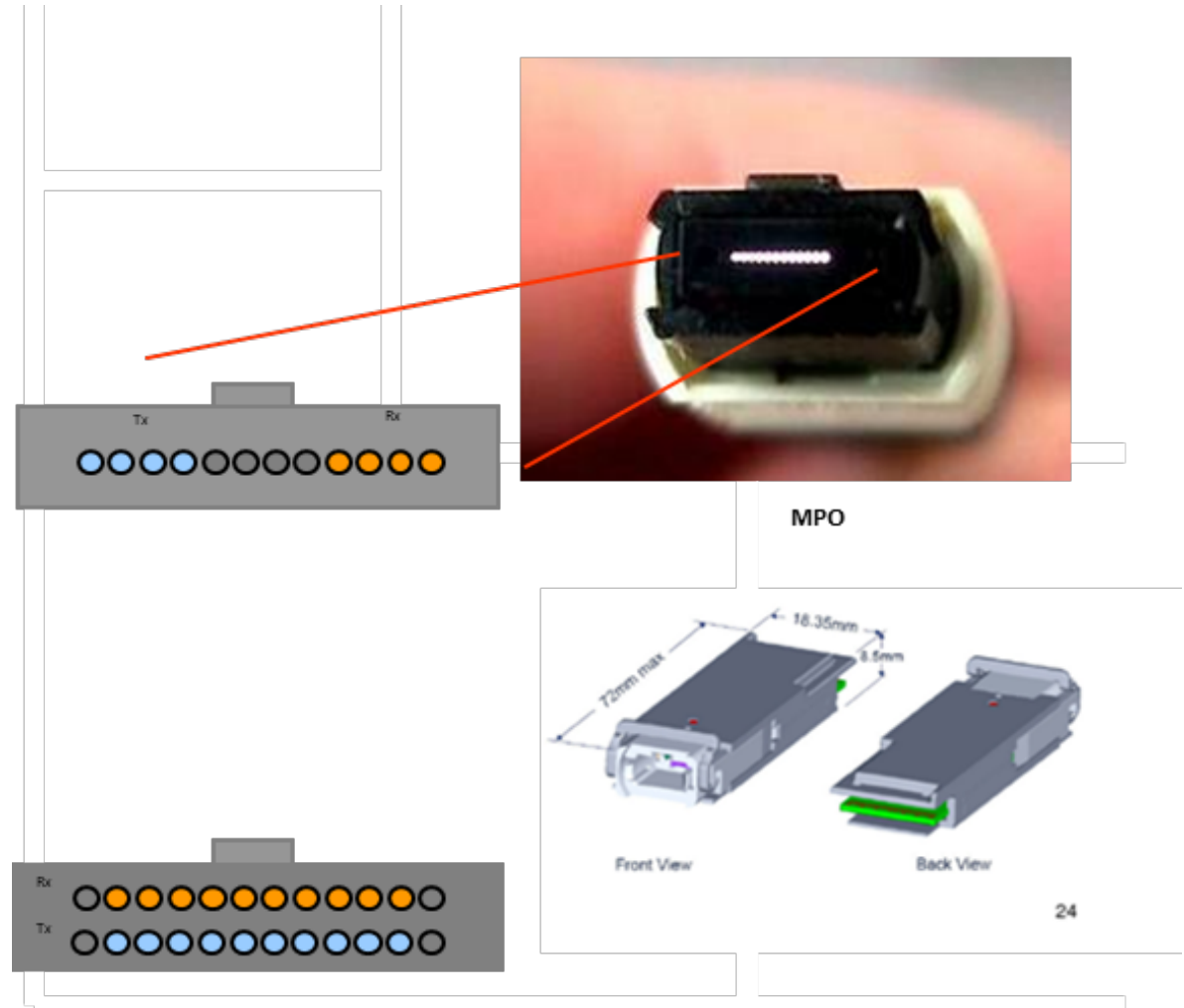
- **40GBASE-SR4**

- Transmisión 4x10G
- Cable 8f OM3

- **100G**

- **100GBASE-SR10**

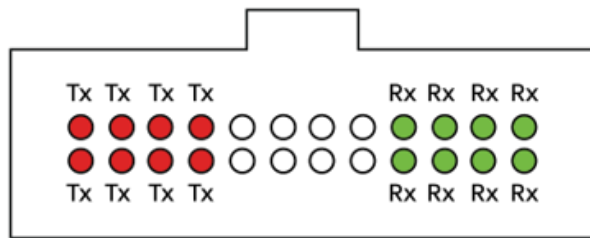
- 10x10G
- Cable 20f OM3



Opción Ethernet 400G-SR8

Soluciones 200G y 400G de IEEE

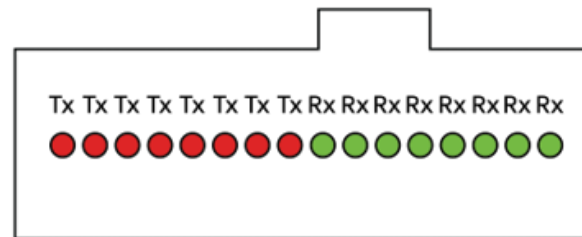
Conector MPO / MTP®



Variante 1

MTP de 24 fibras

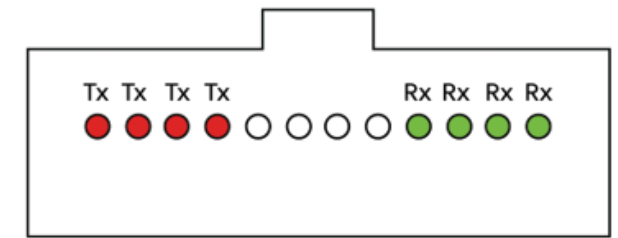
- MTP de 24 fibras ya disponible



Variante 2

MTP de 16 fibras

- Nuevo conector MPO/MTP FILE 389.1
- La versión multimodo es con pulido en ángulo (APC)



Variante 3

MTP de 8 fibras

- Nota: Sol 353.4 400G-SR4.2 y 400G-DR4/XR4
- MTP de 8 fibras fácilmente disponible

Conector MTP y/o MPO Código de colores



Cassettes



24-Fiber MTP
(cassette back view)



12-Fiber MTP
(cassette back view)

Multimodo con bota de color en MTP



Red - 24 Fiber



Aqua - 12 Fiber



Gray - 8 Fiber

Monomodo con bota de color en MTP



Red - 24 Fiber



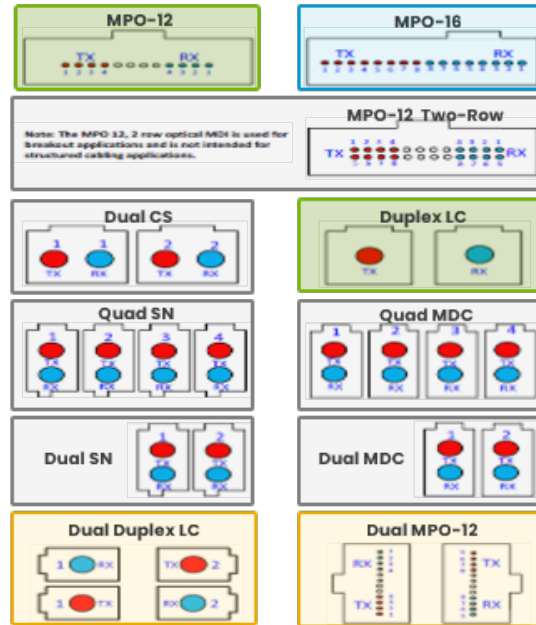
Black - 12 Fiber



Gray - 8 Fiber

Opciones de conectividad de la unidad MSA

QSFP-DD, -DD800 (Sucedendo ahora)



Conectores 16F MTP vs. 16F VSFF



- 400G y conectividad anterior dominada por LC y MPO-12 (FR4, DR4, SR4)

- Las conexiones multimodo se vuelven angulares

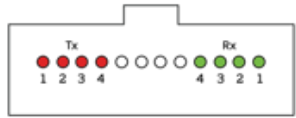
- Estos conectores siguen siendo los más populares a 800G, pero desplegado con 100G / carril, vientre con vientre

- MPO-16 se vuelve más popular a 800G, utilizando 802.3df a 100G / carril

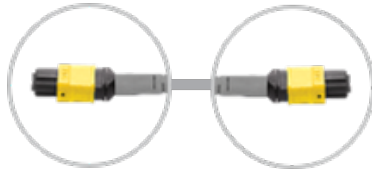
- Existen otras opciones basadas en dúplex y MPO-24 con conectores VSFF, pero en un volumen significativamente menor

Breakouts

Canal 400G Hoy



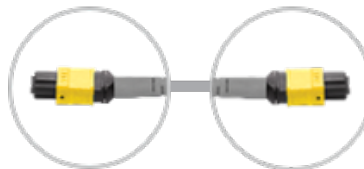
400G-
DR4/XDR4



8F MTP Array Cord



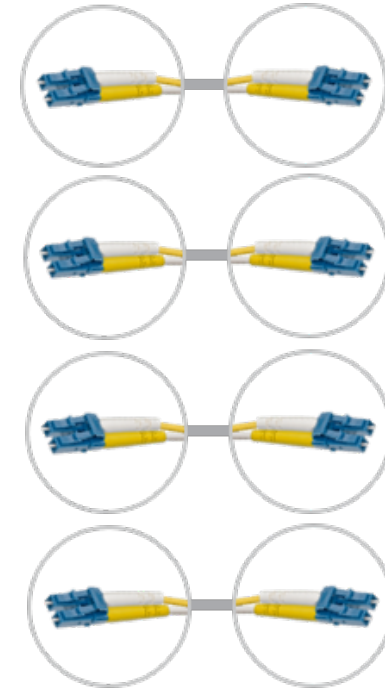
MTP-MTP
Adapters



Base8 MTP Trunk



MTP-LC Cassette



LC-LC Patch
Cords



QSFP28 100G-DR

- Cada línea del transceiver es de 100G
- Convierte canales de 8 fibras en 4 canales dúplex
- Crea 4x100G canales/modulo
- Ventajas de espacio, potencia y coste

Transceivers de 400 Gbps



SFP+
10G



QSFP+
40G



QSFP28
100G



QSFP-DD/OSFP
400G

Arreglo @ 400G: Switch-a-Switch

8-fiber SM Configuración (10/40/100/200/400G) con DR4

OS2

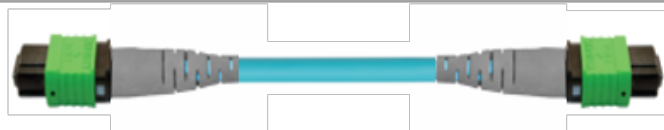


- Cableado central con base de 24 Fibras
- Conexiones Paralelas (8-fibras) al equipo
- Soporta 400G-DR4/XR4 sobre fibra OS2

Conectividad compatible con aplicaciones 400G+

Conjuntos de productos de Leviton para IA y otras redes de alta velocidad

8F Multimodo y Monomodo



Pulido angulado Array Cord

- La interfaz angular es para transceptores de 400G o superiores.
- Matriz de diámetro reducido para conexiones de ultra alta densidad y mejor flujo de aire en el rack.
- Se puede dividir en 1x8F MTP o 2x4F MTP.

8F y 16F MTP placa adaptadora para HDX y E2XHD

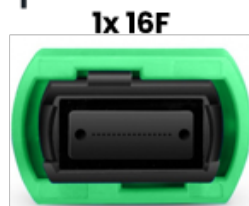


- Conjunto completo de placas acopladoras para dar cabida a la tradición Interfaces 8F y 16F emergentes



16F MTP Array Cords and Arneses

- Una fila de 16 fibras, llave de desplazamiento
- Se aplica como interfaz a los transceptores SR8 (MM) y DR8 (SM) 400G, 800G+
- Agregación y término de fibra hasta APC 2x8F



8F MM/SM APC y 16F MTP

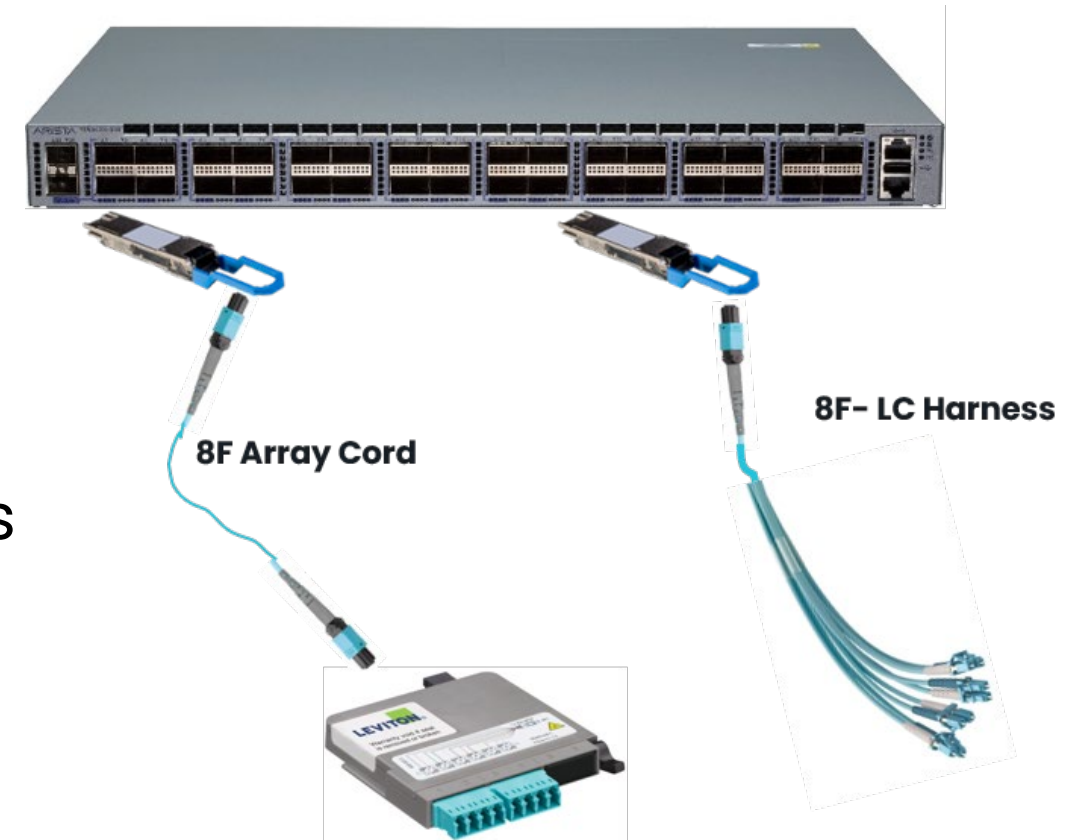
Plenum and CPR B2ca Troncales

- Cableado global y estructurado
Solución para un mayor número de fibras
Consolidación en gastos generales
Bandejas y rejillas



Conversión/Breakout Cassette o Harness

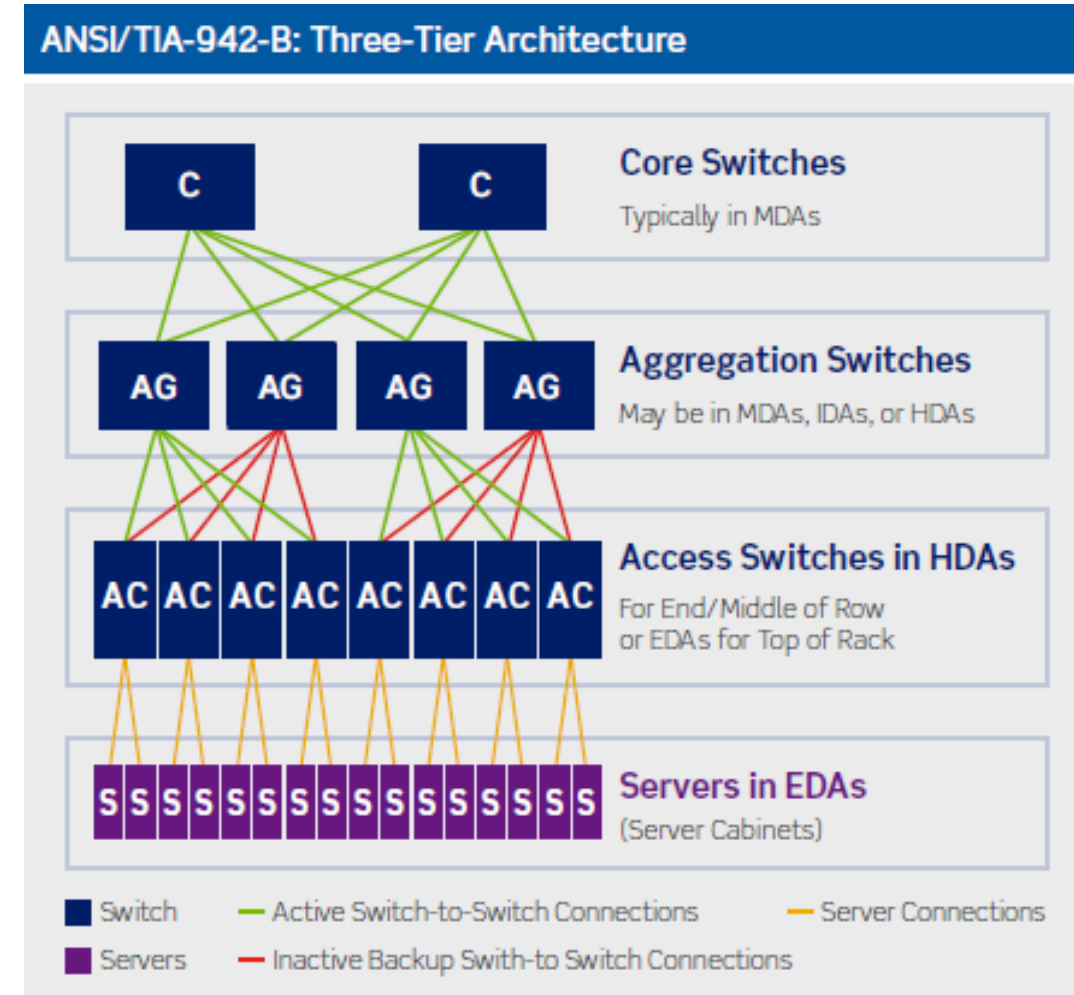
- Segregación de puertos de óptica paralela de 8 fibras (8F) a puertos dúplex individuales
- Con Cassette o Harness
- El cassette permite extender la interfaz de 8 fibras (8F) a ubicaciones remotas mediante el uso de troncales.



Arquitecturas de Data Center

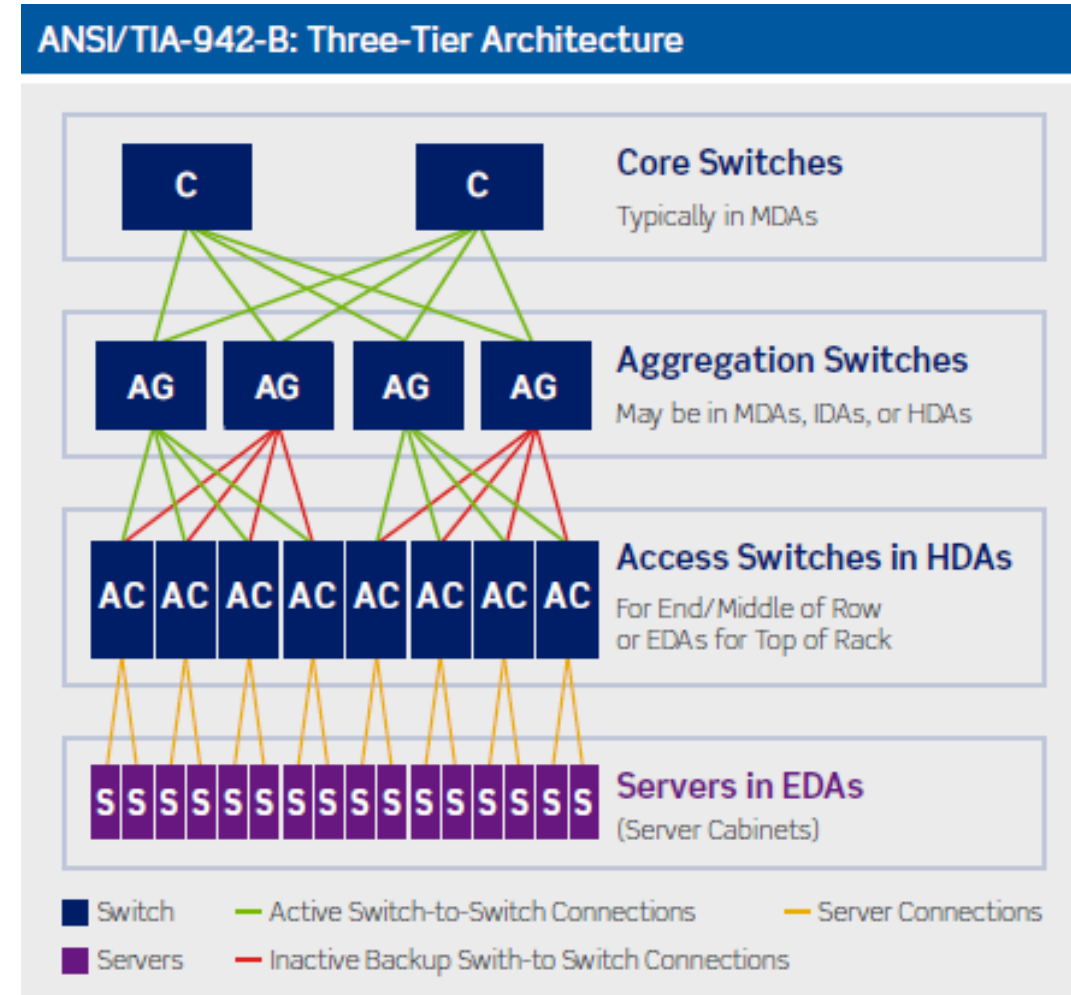
Arquitectura de tres niveles

- Conmutadores de núcleo, agregación y acceso.
- **Se considera anticuada**
- **Todavía habitual en Enterprise**
- **Típicamente oversubscribed**
- **NO optimizada para latencia**



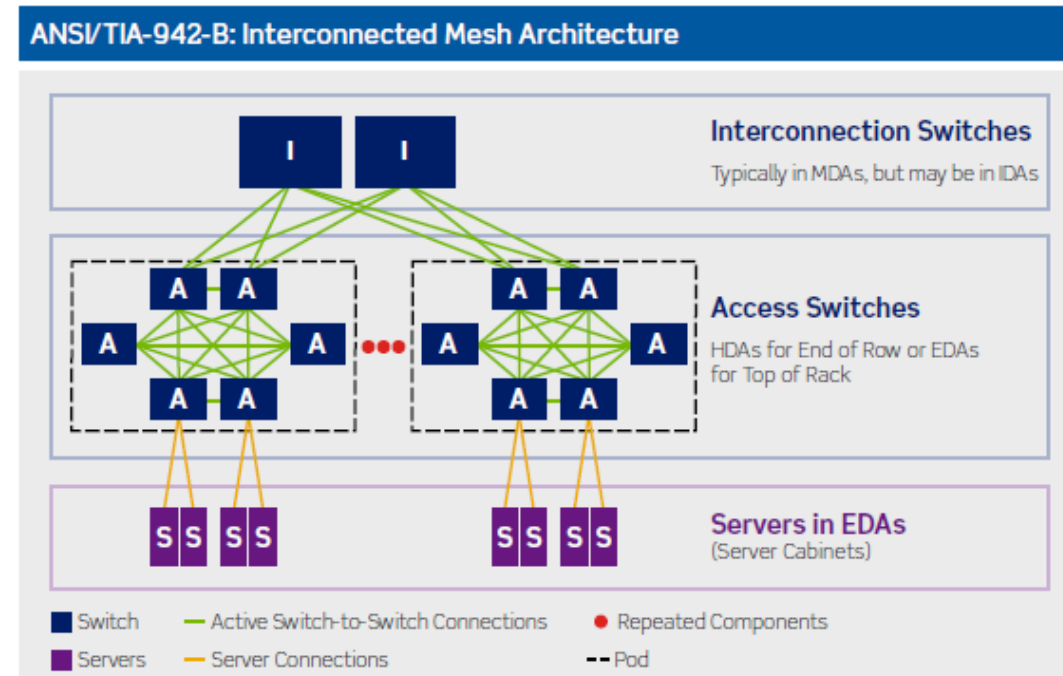
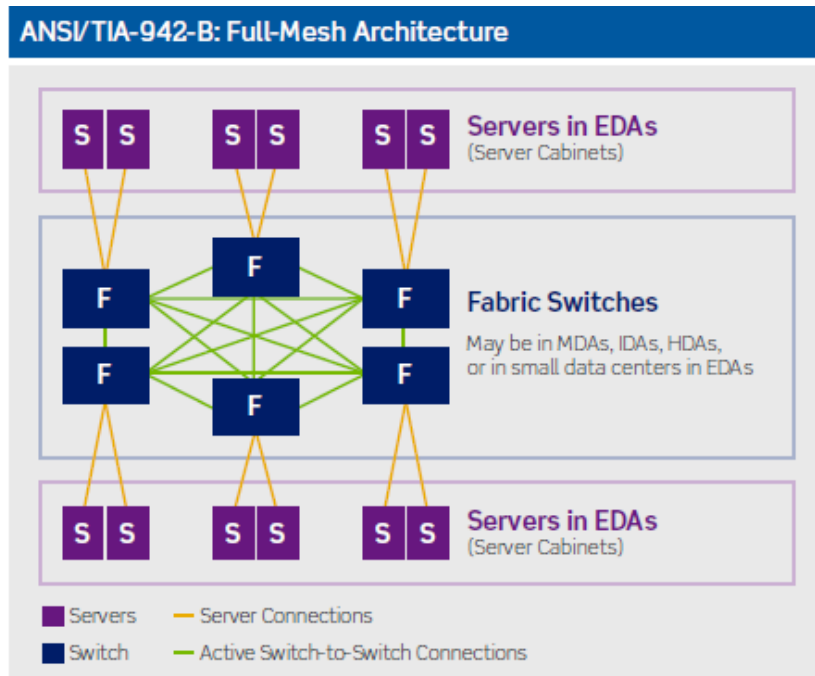
Arquitectura Spine-Leaf

- 2 capas. Cada switch Leaf se conecta a todos y cada uno de los Switches Spine.
- **Necesaria para Cloud/Hyperscale**
- **Minimiza latencia**
- **Típicamente totalmente subscrita**
- **Switch radix importante por el elevado número de conexiones**



Arquitecturas de Malla (Mesh)

- FULL-MESH: Cada switch está conectado a todos los demás
- INTERCONNECTED-MESH: Similar a full-mesh, pero más fácil de escalar



Requisitos del Cableado Mesh

NVIDIA DGX SuperPod Compute Fabric

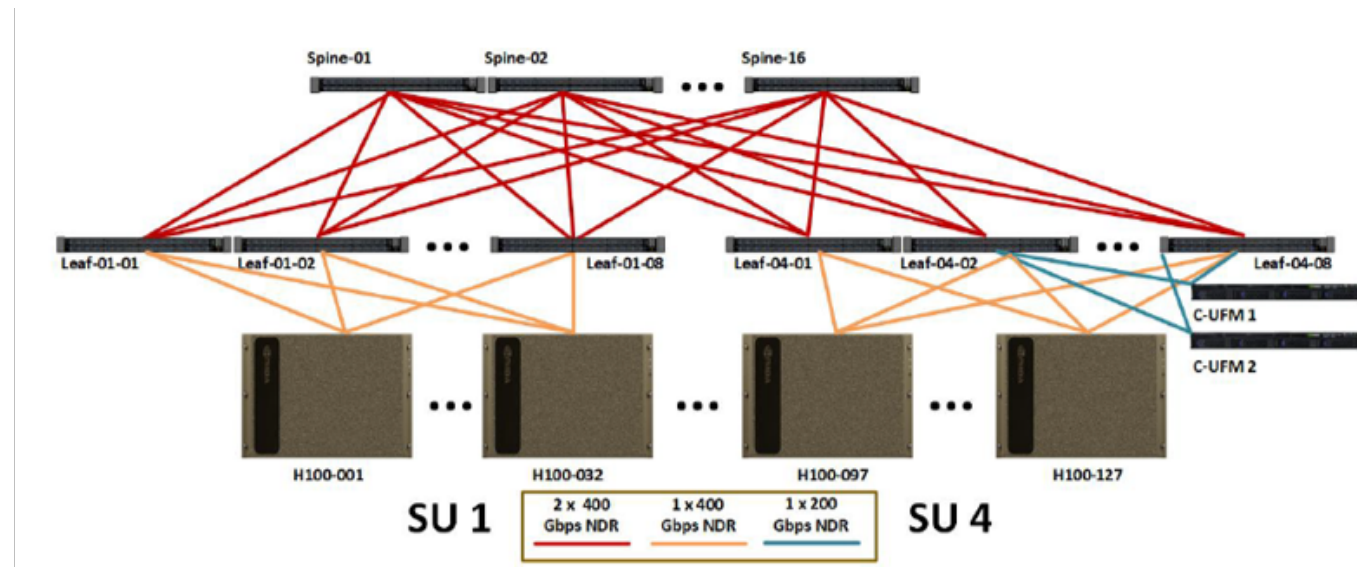
- Red de malla de (400/800Gbs).
- Cada servidor tiene 1 GPU conectada a cada uno de los 8 switches Leaf de 400G.
- Los switches Leaf están conectados a cada switch Spine a 800G.
- Lo típico son cables MPO multimodo (MM) APC, pero pueden ser monomodo (SM) LC.

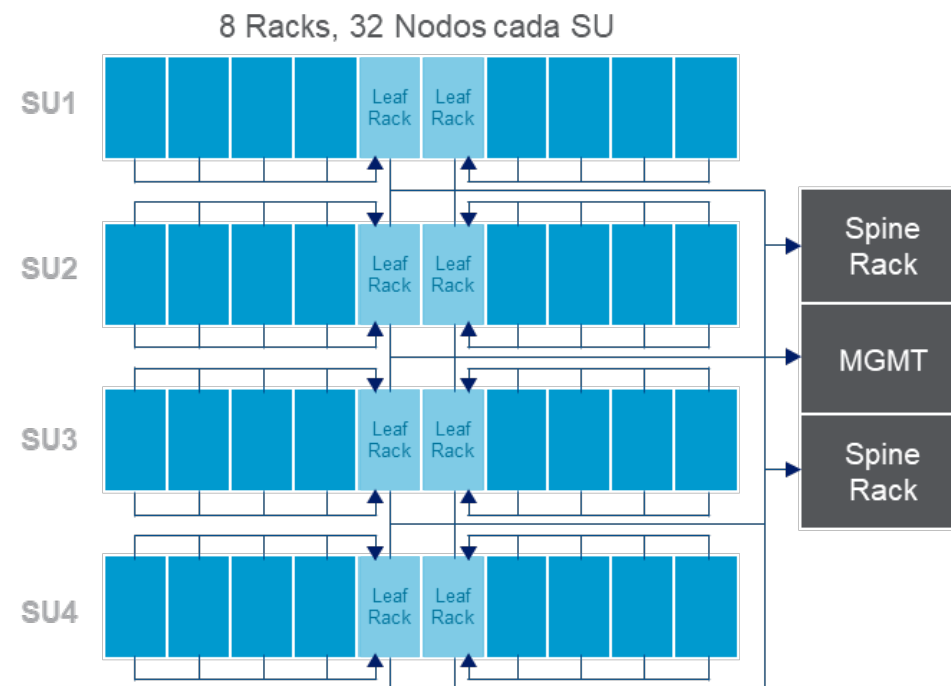
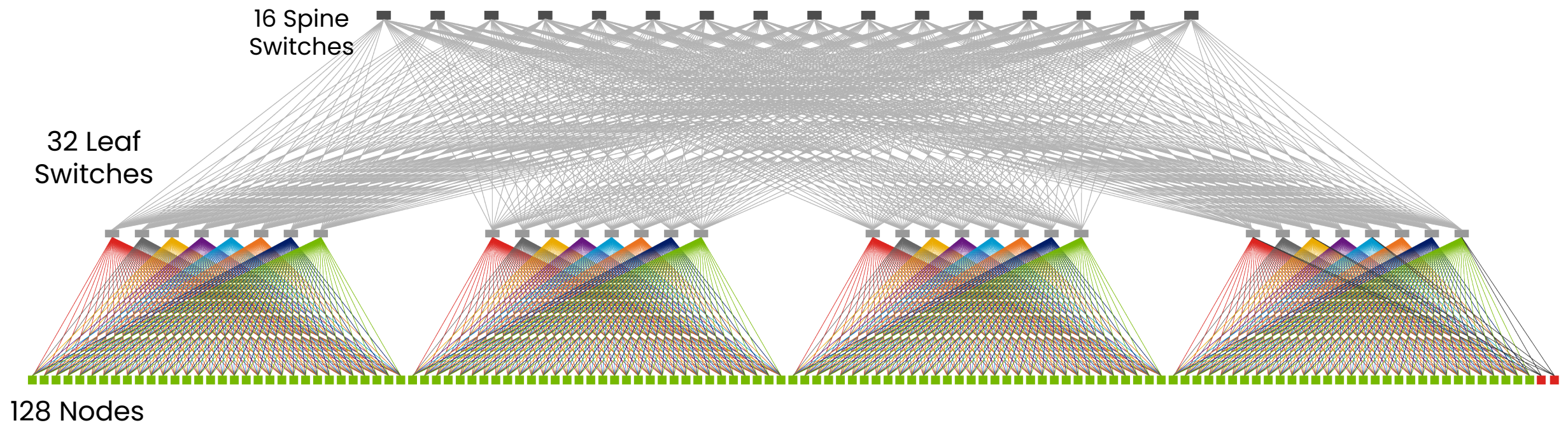


1024 cables



1024 cables

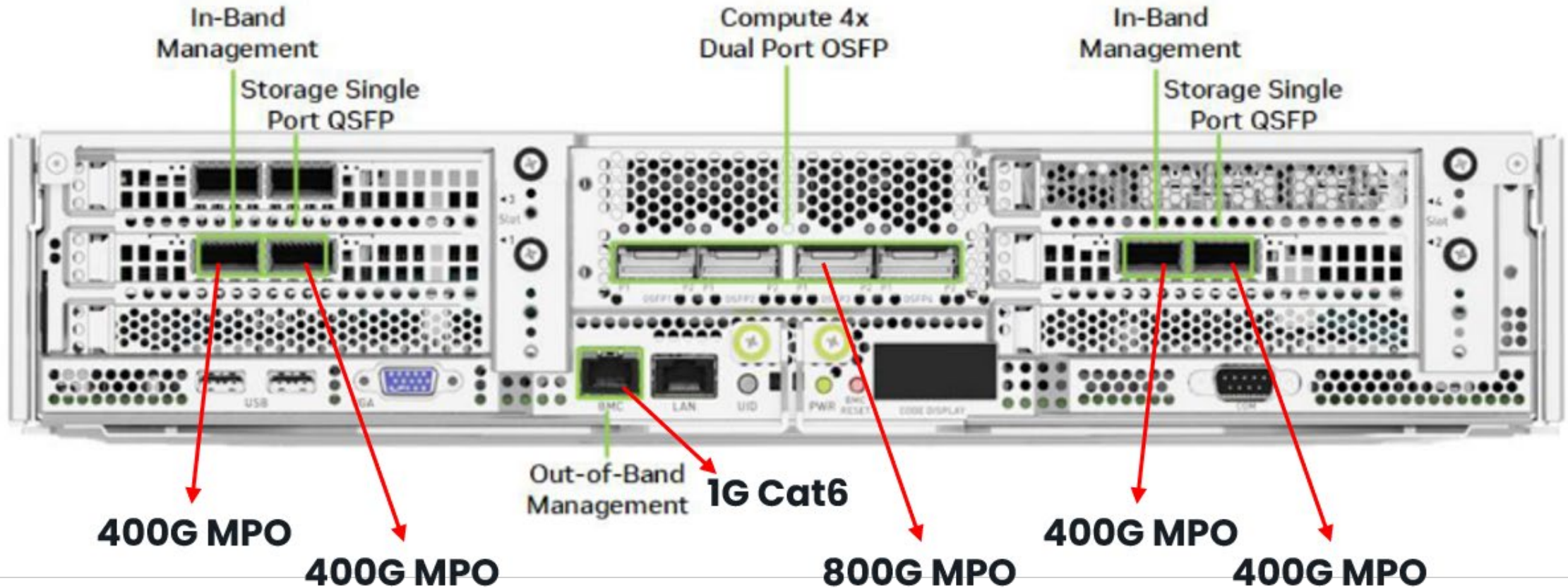




NVIDIA DGX H100 Network Sled Vista

posterior

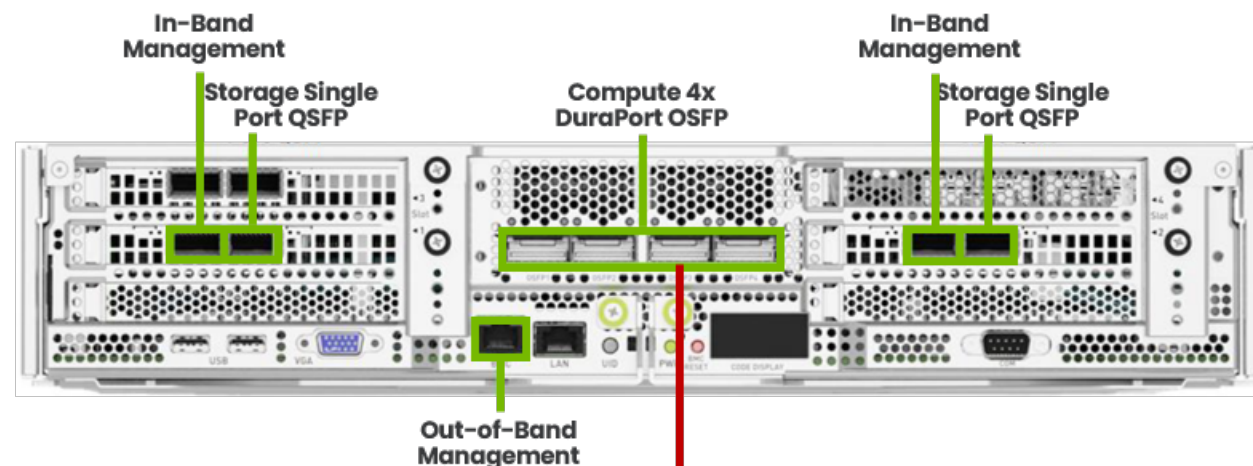
Puertos



NVIDIA DGX H100 Network Sled Vista Posterior

Requisitos de Puertos y Fibras

DGX H100 Network Ports



Nuevo Conector: 8F MM APC



Transceptores de
doble puerto
(2x400 Gbps) para
conexiones de
computación

Requisitos Por Servidor

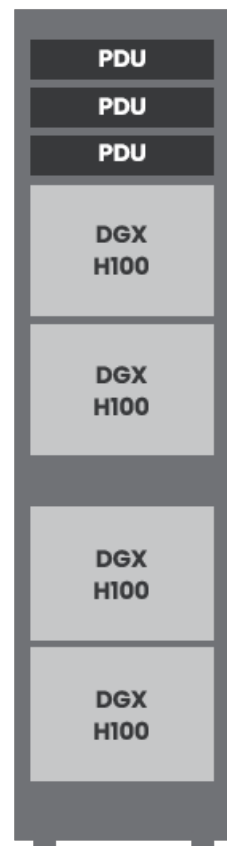
Función	Puertos MPO 8	Fibras
Compute	8	64
Storage	1 o 2	HASTA 16
In-Band	1 o 2	HASTA 16
Total	10 a 12	HASTA 96

Función	Puertos Cobre
OOB Cat 6	1

Ejemplo

Requisitos de cableado Rack de Servidores AI

NVIDIA DGX H100 Server Rack
4 Servidores Por Rack



Requisitos Por Servidor

Función	Puertos MPO 8	Fibras
Compute	8	64
Storage	1 o 2	HASTA 16
In-Band	1 o 2	HASTA 16
Total	10 a 12	HASTA 96

Función	Puertos Cobre
OOB Cat 6	1

x4

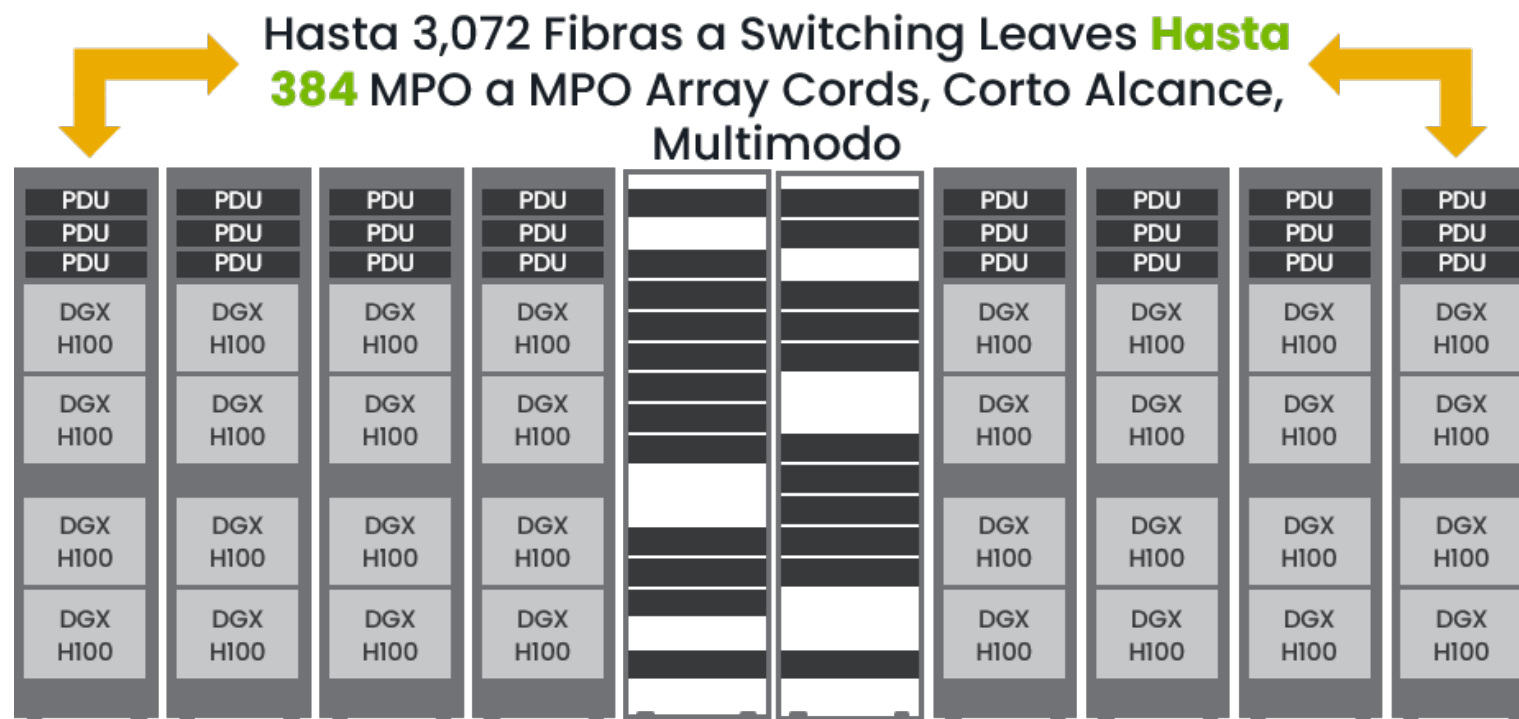
Requisitos Por Rack

Función	Puertos MPO 8	Fibras
Compute	32	256
Storage	4 o 8	HASTA 64
In-Band	4 o 8	HASTA 64
Total	40 a 48	HASTA 384

Función	Puertos Cobre
OOB Cat 6	4

Ejemplo

NVIDIA AI "Scalable Unit" (SU)



Todo conectado a Racks de Switches

Requisitos Por Rack

Función	Puertos MPO 8	Fibras
Compute	32	256
Storage	4 o 8	HASTA 64
In-Band	4 o 8	HASTA 64
Total	40 a 48	HASTA 384
Función	Puertos Cobre	
OOB Cat 6	4	



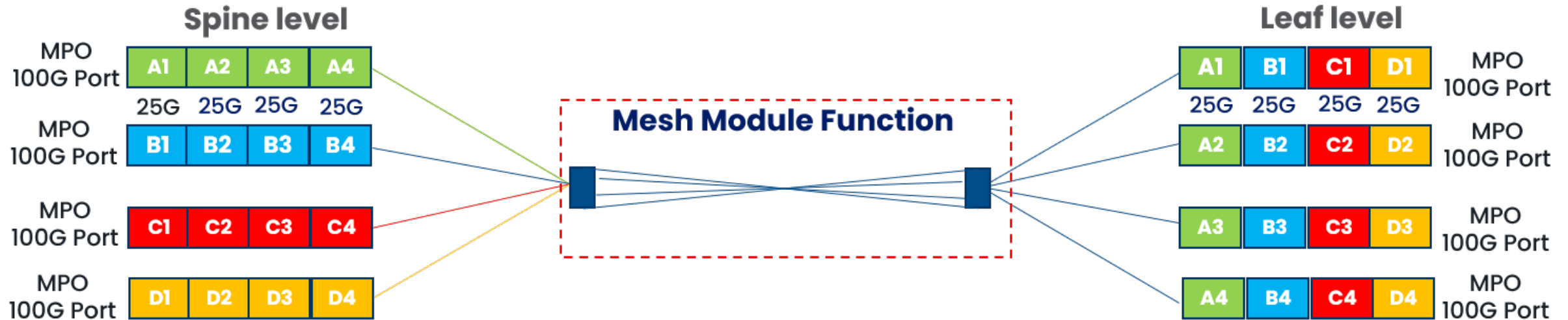
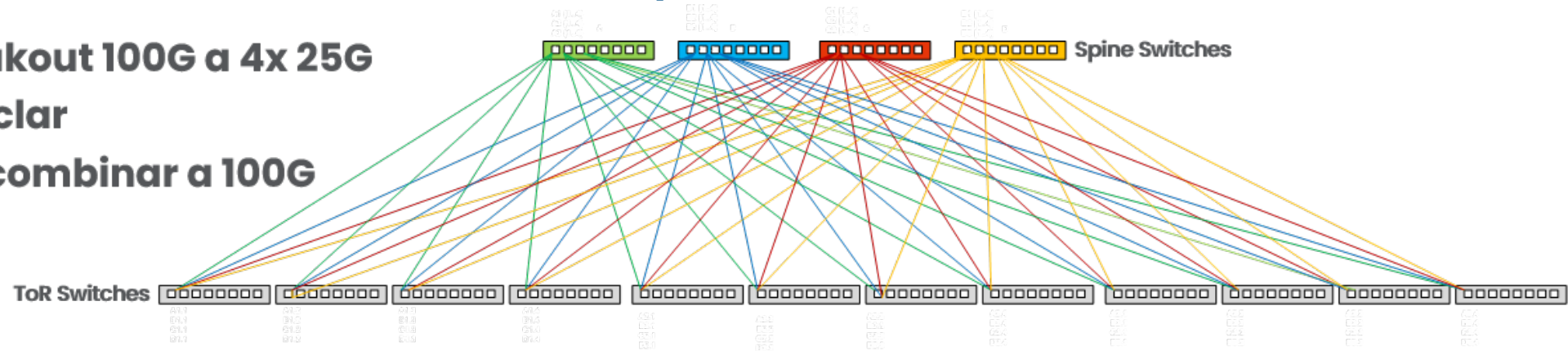
Requisitos Por SU

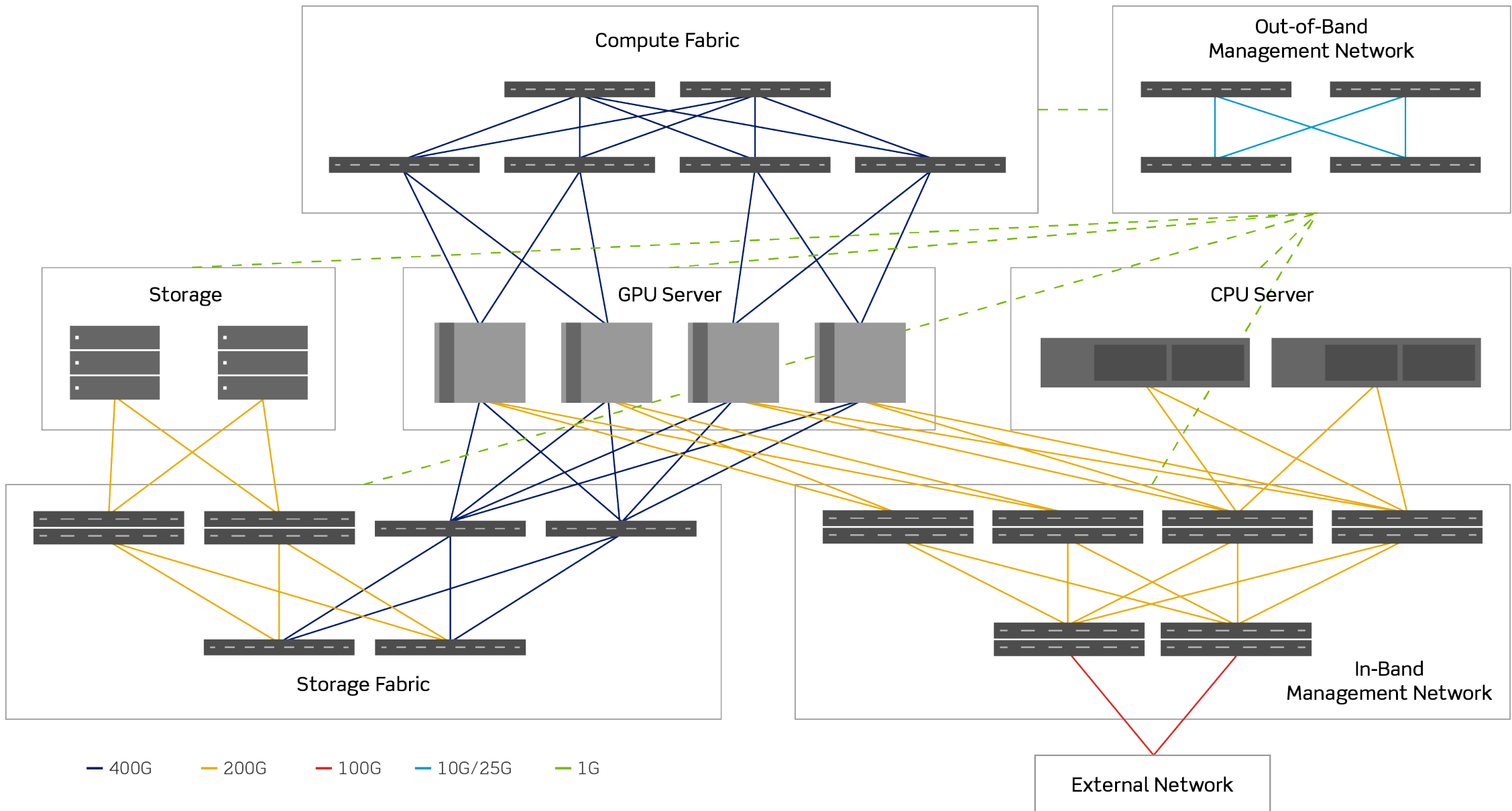
Total Puertos MPO 8	HASTA 384
Total Fibras	HASTA 3,072
Total Enlaces Cobre	32

Función Mesh "Breakout" en Arquitecturas de Data Center

Conexiones Leaf a todos los Spine

- Breakout 100G a 4x 25G
- Mezclar
- Re-combinar a 100G







La **probable** realidad del mallado

En el Mercado existen Data Center Design Team

- Se asocia con los clientes para crear diseños de infraestructura de red eficientes y rentables.
- Expertos con experiencia en equipos activos en el mundo real y conocimientos de planificación.
- Especialización en aplicaciones.
- Modelado anticipado de diseños de red.



Las empresas con los conocimientos para hacer estos proyectos
tienen el futuro garantizado

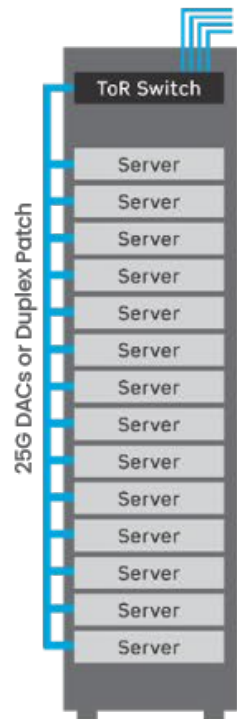


Velocidad de implementación

La velocidad de implementación es un problema

El recuento promedio de fibras está aumentando significativamente ¿Por qué?

Bastidor de servidor tradicional



2 to 6 MPO Uplinks per Rack =
24-Fiber or 48-Fiber
MTP trunks

Rack de servidores
NVIDIA AI de uso
general

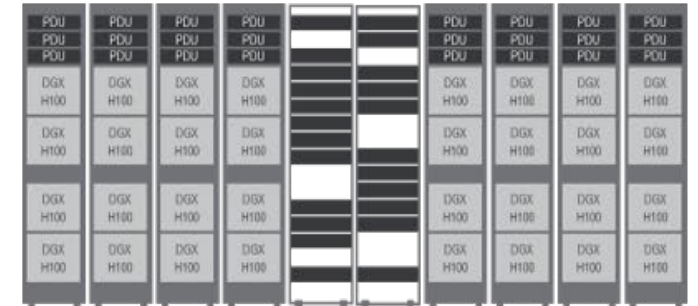


Requisitos de enlace ascendente por rack
~ 10 veces los tradicionales

Function	MPO 8 Ports	Fibers
Compute	32	256
Storage	4 or 8	UP TO 64
In-Band	4 or 8	UP TO 64
Total	40 to 48	UP TO 384

Conclusiones clave:

- Los cables de fibra óptica de mayor capacidad son más comunes que nunca (288F+).
- Esto impulsa la necesidad de conectividad de paneles de mayor densidad (hasta 3456F por unidad de rack).



Una "unidad escalable" requiere 16
troncos de 192F o 2 troncos de 1728F.

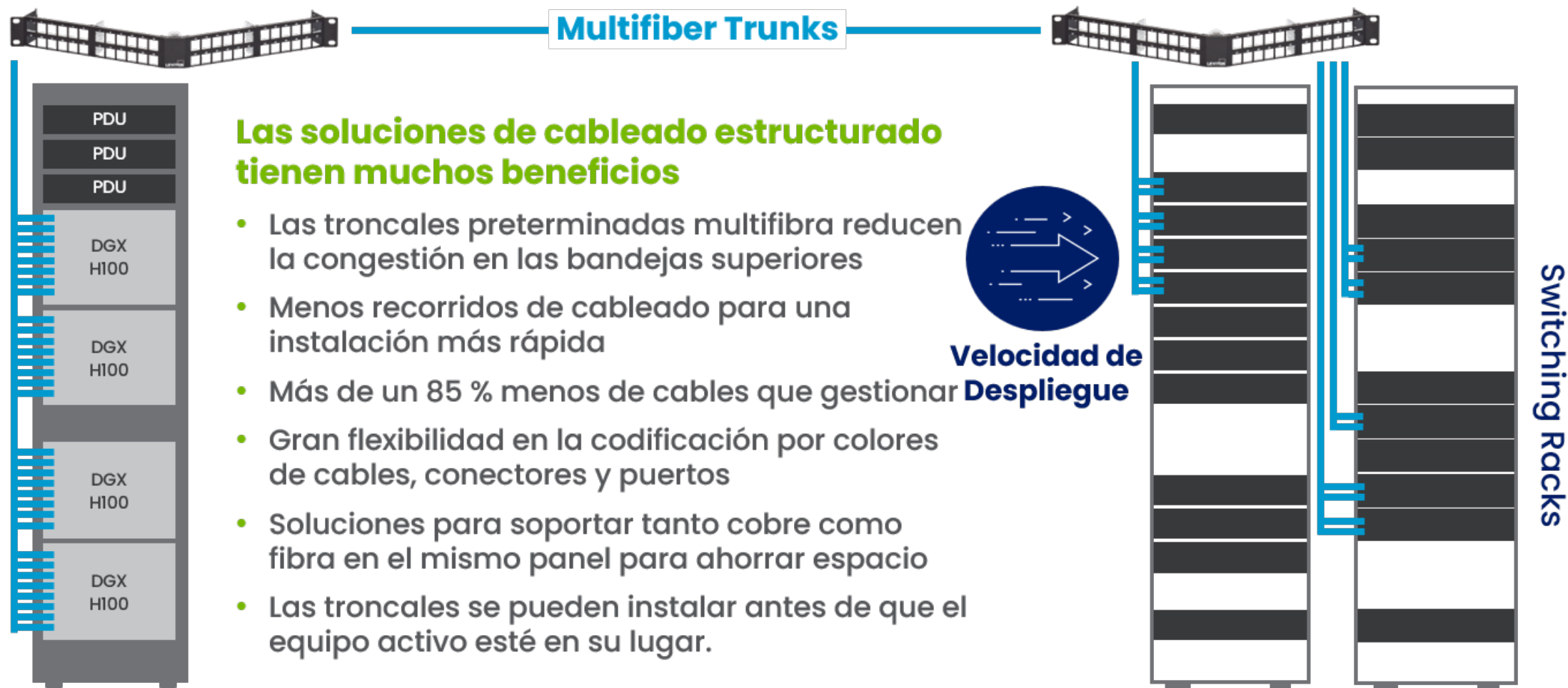


Los clusters de IA (cientos de
racks) crean rápidamente esto



Ejemplo

Cableando la Red IA – Network Solutions



Ejemplo de aplicación

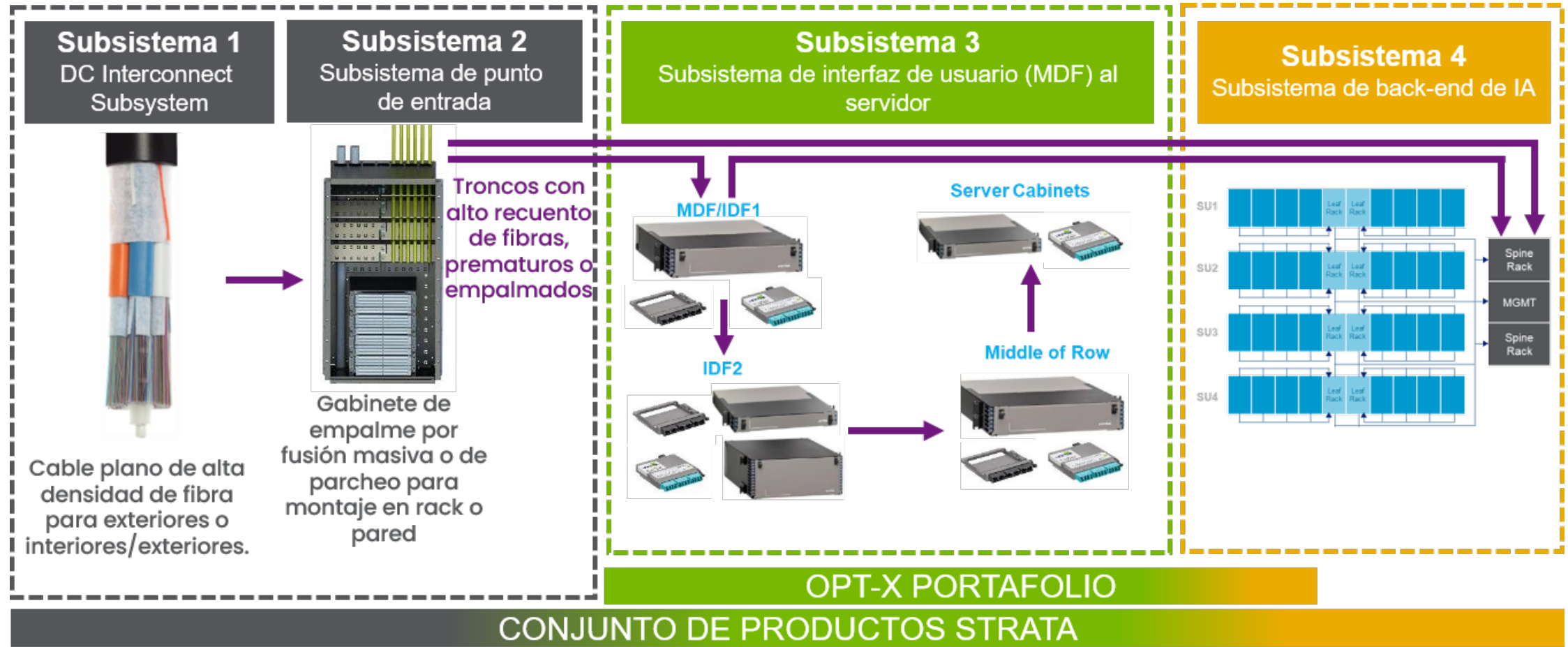


Solución Leviton para Redes AI

Ejemplo de aplicación



Solución integral para Data Centers



Reliable Network Performance

PAST THE OUTER LIMITS

Escenarios comunes de larga distancia

Cámaras de seguridad al aire libre



- Estacionamientos
- Portones de acceso
- Casetas de vigilancia
- Muelles de carga

Ubicaciones al aire libre



- Universidades
- Hospitales
- Entre edificios

Grandes espacios al interior



- Almacenes
- Centros de distribución
- Instalaciones deportivas

Instalaciones Industriales



- Centrales eléctricas
- Canteras
- Fabricación

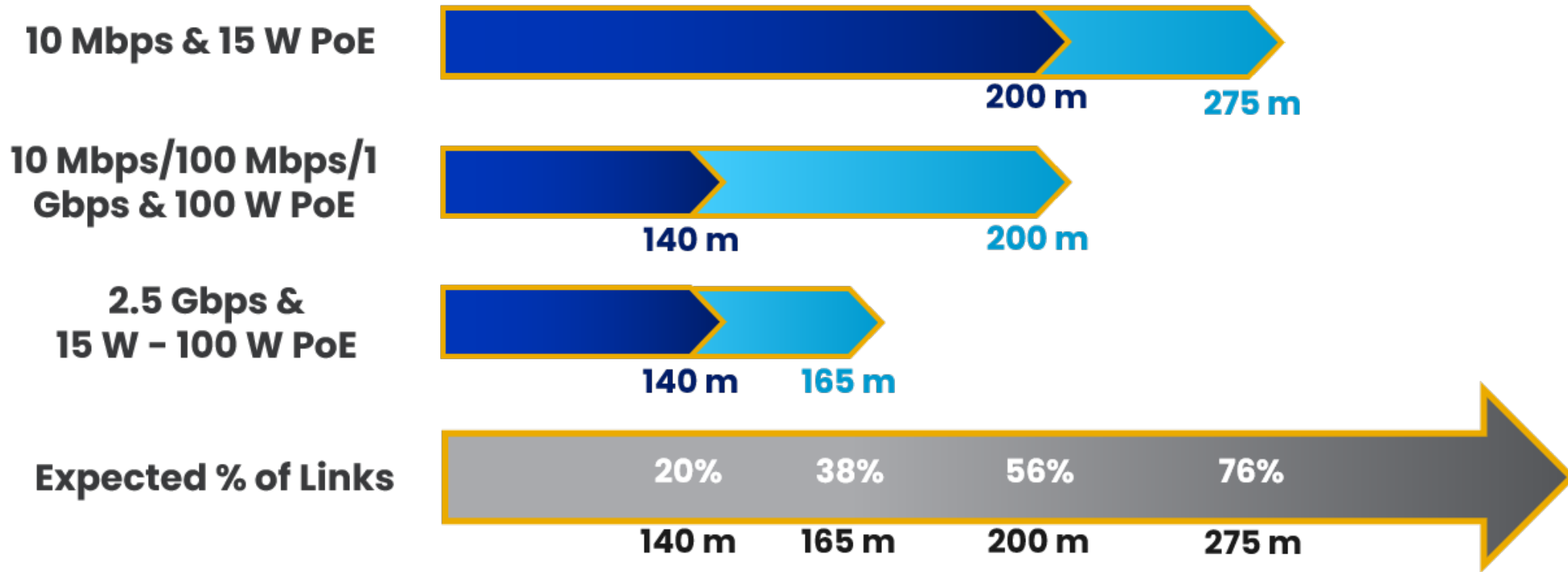
"Aquél...."



- Cámara
- Estación de control de acceso
- Punto de acceso inalámbrico

La mejor cobertura del sector

para los enlaces más largos



Resumen

- > Potencia, refrigeración, ubicación geográfica, latencia y velocidad de implementación son más importantes que nunca en el despliegue de redes de IA
- > Cableado estructurado puede ayudar a abordar estos requisitos críticos
- > Leviton tiene un conjunto completo de productos y sistemas globales definidos para satisfacer las necesidades de las redes de IA en todo el mundo



¿Preguntas?

Antonio Cartagena
Leviton Network Solutions
acartagena@leviton.com
+1 787 903 0304



Leviton.Latam



@leviton.latam



Leviton.Latam



@levitonlatam