

July 13, 2026

NPO Takes the Baton: How Next-gen Optical Interconnect is Now Focused on NPO // Multiple Scale-up and Scale-out projects are now orienting towards NPO.

17 minutes

By Daniel Nishball, Wega Chu, Clara Ee, Gerald Wong, Terence Ong, Julien Martin-Prin and Myron Xie



Executive Summary

- We have discussed the challenges that vendors are facing in quickly bringing CPO to market. NPO and Pluggable CPO provide an alternative that delivers power efficiency while sidestepping the attach yield problem and enabling optical engine vendor diversity.
- For Scale-up, Amazon's Trainium4 will use NPO for both the UALink and NVLink versions; Huawei's Ascend 950 UB-Mesh-Pod will also use NPO.

Nvidia's VRU NVL576 remains on CPO for now, and they could pivot to NPO. They could also opt to scrap the optical cross-rack solution and try again with Feynman.

- We expect the total key component TAM for CPO and NPO to reach \$31.5B by 2030. NPO should drive most of the growth in this TAM in 2028 and 2029, representing 60-70% of optical engines shipped in those years. By 2030, however, CPO will likely represent the majority of OE shipments as more platforms successfully develop their CPO solutions and as OCI-MSA becomes an option.

System Vendors Falling back to NPO

We have [discussed at length the various challenges](#) facing even the lower risk and simpler use cases for CPO, namely scale-out switching. Nvidia has been struggling with its CPO optical engine (OE) production and with fiber attach issues and is currently hamstrung by lower than ideal Optical Engine attach yields. This might not only stymie their scale-out switch projects but could also derail plans to ship their VRU NVL576 8-rack scale-up world size system if these issues persist. Ramping scale-up optics is not easy and carries significant production risk.

Meanwhile, the first systems to use optics for scale-up appear to be tapping NPO and Pluggable CPO first instead of CPO in part to de-risk, while still harnessing some of the benefits of CPO. Amazon's Trainium 4 will be the first large custom silicon system to attempt adopting NPO, while Huawei's Atlas UB-Mesh-Pod built around the Ascend 950 is also adopting NPO for scale-up networking.

On the scale-out front, Meta is playing many angles. Not only does it have a DR Optics CPO scale-out switch program built by Celestica, but it is also exploring an NPO scale-out switch program. Meta is also driving an OCI-MSA based CPO scale-up Tomahawk 6 based switch program that could be used for its own MTIA system or in conjunction with systems from merchant

silicon vendors that are also part of the OCI MSA, but as deployment is still a few years out – Meta still has time to wait for the kinks to be ironed out.

The implication is that NPO will drive much of Optical Engine shipments in the early days of CPO/NPO – making up most shipments in 2028 and 2029 before scale-up CPO matures in earnest in 2030.

CPO and NPO TAM Build-Up (Install Base Convention)						
	Units	2026	2027	2028	2029	2030
NPO Scale-Out 3.2T Equiv OEs	Units	0	484,704	1,341,360	1,347,581	366,638
NPO Scale-Up 3.2T Equiv OEs	Units	399,000	1,312,500	13,430,552	40,585,385	24,841,385
NPO 3.2T Equiv OEs	Units	399,000	1,797,204	14,771,912	41,932,966	25,208,024
CPO Scale-Out 3.2T Equiv OEs	Units	11,847	1,407,865	4,776,733	6,362,874	9,081,450
CPO Scale-Up 3.2T Equiv OEs	Units	0	0	2,731,308	15,182,183	46,810,600
CPO 3.2T Equiv OEs	Units	11,847	1,407,865	7,508,040	21,545,057	55,892,050
Total 3.2T Equivalent Optical Engines - CPO + NPO	Units	410,847	3,205,069	22,279,953	63,478,023	81,100,074

Source: [SemiAnalysis AI Networking Model](#)

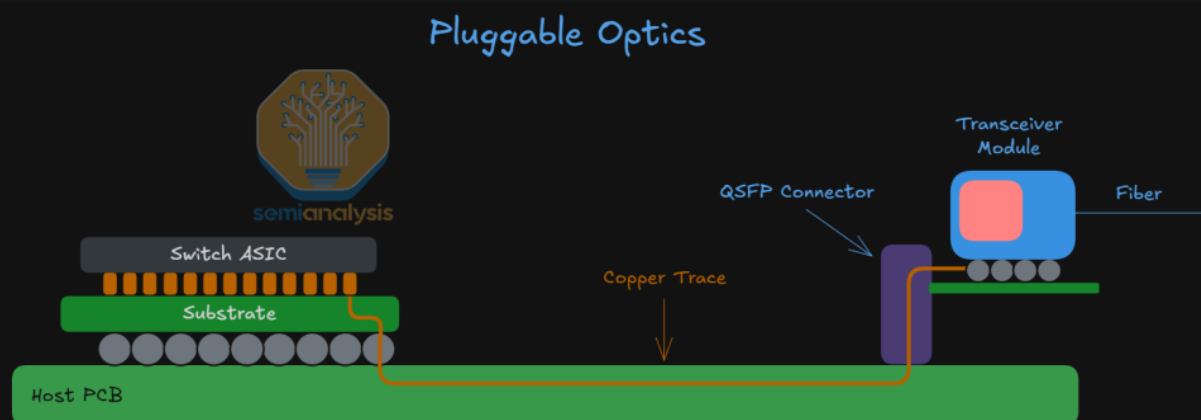
Recall that in the AI Networking Model, estimates and outputs such as pivot tables are on an install base convention. The **“Install Base Convention”** allocates units and \$ spend to the quarter in which the cluster is brought online. Each component and company will differ in lead times and revenue recognition and will require a different lead time to reconcile to revenue. Our company reconciliation tabs will use a lead time that is typically a “bring forward” of the unit demand indicated under the install base convention to reflect that fact that equipment is usually purchased in advance of cluster installation.

Pluggables vs NPO vs CPO vs XPO and CPC

To understand the current battleground with respect to optical connectivity solutions, let’s quickly recap the definitions and basic architecture of these solutions.

The most common form factor today is the pluggable transceiver – the optical engine sits at the faceplate in a standardized cage, and the electrical signal must travel from the ASIC through the package BGA and across roughly 300mm of host PCB trace before it reaches the module’s connector.

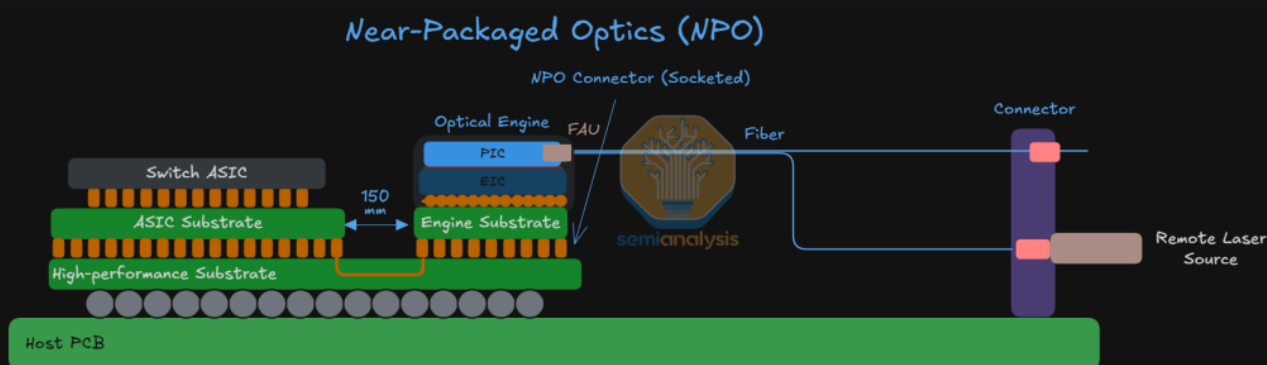
The copper trace is lossy enough that the transceiver module carries a full/partial DSP to retime the signal – and that DSP is most of what the industry is trying to remove since it accounts for roughly half the module's power.



Source: [SemiAnalysis AI Networking Model](#)

A few form factors to replace pluggable optics exist. Near-Packaged Optics (NPO) moves the optical engine of the faceplate and lands it next to the ASIC package – but critically, not on the ASIC substrate. As the diagram below shows, the engine sits on its own substrate, mated through a socketed NPO connector, with the high-speed channel running roughly 150mm through a high-performance substrate rather than an ordinary host PCB – only low-speed signals and power are delivered through the channel.

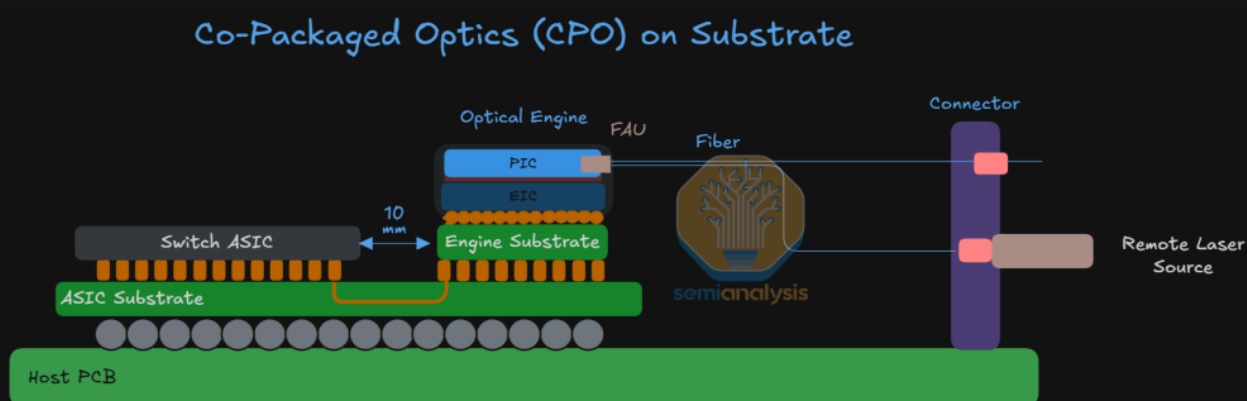
We expect current NPO implementations to adopt 200G per lane PAM4 modulation and that most NPO solutions will be a pluggable form factor (as detailed by the CPX-MSA), which enables vendor diversity within the rack.



Source: [SemiAnalysis AI Networking Model](#)

Co-Packaged Optics (CPO) puts the engine on the ASIC substrate itself, cutting the electrical channel to ~10mm. The key element at play here is that CPO is on the ASIC substrate – this is what differentiates CPO vs NPO as opposed to whether the optical engine is pluggable. For example, Nvidia's Quantum CPO solution uses three 1.6T optical engines on an assembly that is pluggable onto the substrate – and for most, it still counts as "CPO".

To clarify what we mean, if the optical engine is pluggable AND it connects to the ASIC substrate, then we will call it "Pluggable CPO", but if it is permanently attached to the ASIC substrate, then we will simply refer to this as "CPO". A few industry players have set out to push for CPO solutions to be pluggable as this is the only way to open the ecosystem to third-party OE vendors and as it also Pluggable CPO dramatically more deployable by de-risking yield and rework.



Source: [SemiAnalysis AI Networking Model](#)

What is the impetus for exploring CPO in the first place? CPO saves power because it eliminates the DSP by placing the optical engine next to the XPU or switch ASIC. However, if this is not matched by a commensurate downgrade in power for the Switch or XPU SerDes from LR SerDes to shorter reach SerDes, power benefits are limited beyond what can be achieved by Linear Pluggable Optics.

Today's CPO implementations focus on DR Optics based Optical Engines, but if the OCI-MSA version of CPO is achieved, power savings on the optical engine will be even more significant with the optical engine running on 50G NRZ lanes instead of 200G PAM4 SerDes lanes, taking down the OE + ELSFP power down from 5-6pJ/bit in the case of DR Optics CPO down to below 3

pJ/bit for OCI-MSA CPO. Further system power savings can be achieved by eliminating SerDes on the Switch or XPU and using a chip-to-chip interconnect to the optical engine.

CPO also has the benefit of supporting bandwidth density scaling beyond what is possible with OSFP pluggable optics. CPO eliminates the faceplate thermal and throughput limit, since unlike a pluggable module, data transmission occurs within the package. In fact, Nvidia's scale-out CPO switches will exclusively use denser MMC connectors for the faceplate, abandoning the ubiquitous MPO connector form factor. Optical Engines on Interposer will also unlock more overall bandwidth escape from the package as shoreline density can be far greater if OEs are connected directly to the Switch or XPU, much like HBM is connected today.

The technical challenges with implementing CPO from the get-go are the reliability and serviceability issues associated with packaging optical engines close to the ASIC as well as [yield issues at the initial stages of CPO implementation](#) as the supply chain ramps to support growing volumes. On top of this, one of the largest hesitations over CPO adoption stems from the fact that CPO implementations that are not pluggable result in vendor lock in, and the resulting loss of sourcing flexibility, field serviceability and have the consequence of the switch vendor likely taking a 70% gross margin stack on the optical engines vs the 40% gross margin that a third party pluggable Optical Engine vendor might charge.

Meanwhile – NPO and Pluggable CPO offers much of the above benefits but without the vendor lock in. It also inherently sidesteps the attach yield issue that is the current headache for scale-out CPO switch production.

We compare key interconnect solutions from power efficiency and cost per bit perspective in the table below. NPO, CPO, and Pluggable CPO all have the potential to drive better energy efficiency than LPO if all opportunities for power savings are harnessed. Only OCI-MSA on interposer has the potential to deliver even greater power savings, but this will require customers to potentially accept vendor lock in, or at the very least accept that the switch

or ASIC vendor will be the protagonist when it comes to deciding which vendor will be chosen for Optical Engines.

Optical Interconnect Solutions Comparison							
	Traditional Pluggable Module	LRO with Pluggable Switch	LPO + CPC using XPO	NPO	CPO	Pluggable CPO	CPO (OCI-MSA)
Signal Path from Switch to OE	BGA+PCB	BGA+PCB	CPC + Flyover	Substrate + PCB	Substrate	Substrate	Interposer
Optical Engine (OE) Position	Faceplate	Faceplate	Faceplate	On PCB, Pluggable	On Substrate, Not Pluggable	On Substrate, Pluggable	On Interposer
Signal Conditioning?	Fully Retimed	Retimed on Tx Channel	Linear	Linear	Linear	Linear	Linear
Electrical Channel Length	30cm over PCB Trace	30cm over PCB Trace	30cm over Flyover Cables	~5cm over PCB Trace	~5cm	~5cm	~1-2cm
Energy Efficiency - OE + ELS	~13 pJ/bit	~10 pJ/bit	~6 pJ/bit	~6 pJ/bit	~6 pJ/bit	~6 pJ/bit	~3 pJ/bit
Energy Efficiency - Scale-Out + Switch	~27 pJ/bit	~24 pJ/bit	~20 pJ/bit	~20 pJ/bit	~19 pJ/bit	~19 pJ/bit	~15 pJ/bit
Hot Swappable	Yes	Yes	Yes	No	No	No	No
Serviceability	Hot-Swappable	Hot-Swappable	Hot-Swappable	Field-Replaceable Switch must be taken offline	Not Serviceable	Not Serviceable	Not Serviceable
Expected Deployment Timeline	Now	2026-2027	2027 and beyond	2027-2028	2027?	2027?	2030?
Pluggable Option?	N/A	N/A	N/A	Yes	?	?	No
Optical Engine Vendors	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Bundled with Switch or GPU ASIC. Switch/ASIC vendor decides OE supplier	Bundled with Switch or GPU ASIC. Switch/ASIC vendor decides OE supplier	Bundled with Switch or GPU ASIC. Switch/ASIC vendor decides OE supplier
ELSFP Vendors	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers	Third-Party Pluggable Manufacturers

[†] Numbers presented are preliminary and are subject to change.

Source: [SemiAnalysis AI Networking Model](#)

We outline below how total power could look like when contemplating the use of CPO/NPO for scale-out switching. Here we see that NPO, CPO and Pluggable CPO can offer the best energy efficiency, especially if Switch ASICs can be designed with lower power SerDes.

Turning to total system cost from a scale-out perspective, we see that the CPO and CPO (OCI-MSA) have higher total switch selling prices than Pluggable CPO and NPO. This is because we are assuming a 70% vendor margin stack on the Optical Engines and FAU Assembly.

As mentioned above, OCI-MSA based CPO remains the most efficient interconnect with 2-3 times better energy efficient than NPO and 5 times better efficiency than LRO, but as with CPO, it inherently means vendor lock-in.

Optical Interconnect Solutions Comparison								
	Units	Traditional Pluggable Module	LRO with Pluggable Switch	LPO + CPC using XPO	NPO	CPO	Pluggable CPO	CPO (OCI-MSA)
Optical Engines and FAU Assembly	W	0	0	0	480	480	480	179
ELSFP	W	0	0	0	128	128	128	128
Pluggable Modules	W	1,344	1,056	640	0	0	0	0
102.4T Switch Chip + SW Box Items	W	1,420	1,420	1,420	1,420	1,320	1,320	1,245
Total Switch Power	W	2,764	2,476	2,060	2,028	1,928	1,928	1,552
Bandwidth	Tbit/s	102.4	102.4	102.4	102.4	102.4	102.4	102.4
pJ/bit (including Switch)	pJ/bit	27.0	24.2	20.1	19.8	18.8	18.8	15.2
pJ/bit (OE+ELSFP Only)	pJ/bit	13.1	10.3	6.3	5.9	5.9	5.9	3.0
Total Cost of Ownership								
Optical Engines and FAU Assembly	\$	0	0	0	18,393	33,720	16,860	33,720
ELSFP	\$	0	0	0	6,400	6,400	6,400	6,400
Pluggable Modules	\$	51,840	46,080	36,480	0	0	0	0
102.4T Switch Chip + SW Box Items ¹	\$	38,714	38,714	38,714	38,714	38,714	38,714	38,714
Total Selling Price	\$	90,554	84,794	75,194	63,507	78,834	61,974	78,834

¹ Assuming 30% Gross Margin and 3 year warranty

Source: [SemiAnalysis AI Networking Model](#)

What is clear from the table is that the intersection of best power efficiency and lowest cost today is in pluggable CPO and NPO, and this is why these solutions are becoming the focal point for current system deployment. What is not explicitly spelled out in this table is the fact that a less-than-ideal optical engine attach yield tends to increase unit optical engine costs, driving up overall costs for everyone.

Attached Optical Engine Yield Worksheet						
	Yielded Die per Wafer	Yield Rate	Wafers per Month	Yielded Dies/mth	Cumulative Yield	102.4T Equiv Switches/Yr
CPO						
Gross Yielded COUPE Wafers	788	100%	10,000	7,880,000	100%	2,955,000
OE Unit Test Yield	481	61%	10,000	4,810,000	61%	1,803,750
OE Attach Yield per OE		98%	10,000			
Switch Die Attach Yield - 32 OEs	252	52%	10,000	2,519,878	32%	944,954
NPO						
Gross Yielded COUPE Wafers	788	100%	10,000	7,880,000	100%	2,955,000
OE Unit Test Yield	481	61%	10,000	4,810,000	61%	1,803,750
OE Attach Yield per OE		100%	10,000			
Switch Die Attach Yield - 32 OEs	481	100%	10,000	4,810,000	61%	1,803,750

Source: [SemiAnalysis AI Networking Model](#)

As promising as CPO is, it remains risky to leapfrog straight to CPO, and NPO is becoming a viable choice for go-to market in late 2027 and beyond for both scale-out and scale up. Some scale-up programs have elected to start with NPO – for instance Trainium4, while other scale-up programs like Nvidia's VRU NVL576 appear to be pivoting into NPO given challenges with yielding CPO.

Scale-up NPO Projects

Amazon Trainium4

The design for Trainium4 remains in flux still – but as a base case scenario we see Trainium4 as being designed to implement a large world size architecture using both copper backplanes and NPO.

We use the Trainium3 liquid cooled NL72×2 switched (i.e. Trainium3 MAX) as a starting point but this time rather than forming a two-rack world size through cross-rack AECs, we will form a higher density rack and a larger world size using NPO to connect local GPUs to distant switches. Each Trainium4 rack will have 144 GPUs, four GPUs per each of 36 compute trays. For the UALink version, 36× 57.6T UALink switches will be used, with four switches per each of 9 compute trays.

Trainium4 Rack Assumptions		
	Units	Quantity
Compute Trays	#	36
Switch Trays	#	9
Trainium4 per Compute Tray	#	4
Total Trainium4 per rack	#	144

Source: [SemiAnalysis AI Networking Model](#)

A copper backplane will be used to connect each of the 144 Trainiums in each rack locally to every switch within the rack, and NPO will be used to connect each GPU directly to distant switches within each of 7 other racks within the world size. We expect that Amazon could continue to make use of smaller scale-up switches within compute trays to provide additional scale-up switching capacity.

Trainium 4 – NeuronLinkv5 Protocol (UALink or NVLink)

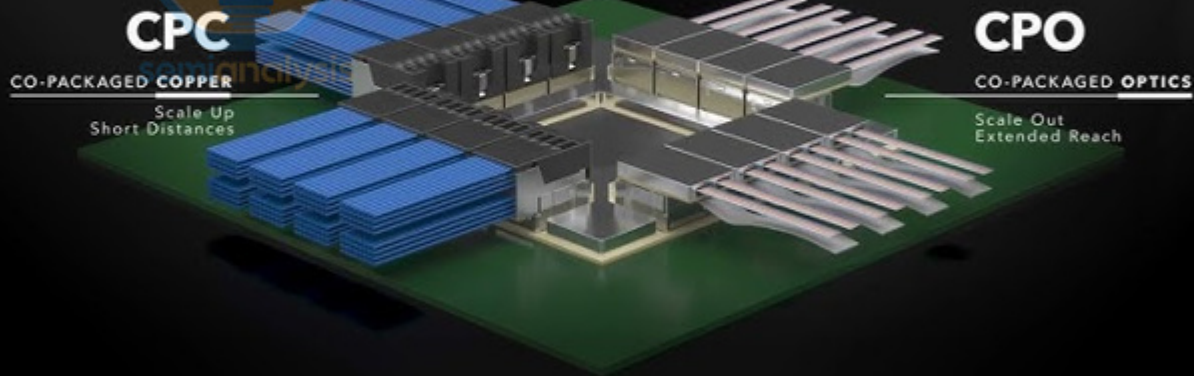
AWS's Trainium scale up interconnect protocol is known as NeuronLink, a modified version of PCIe protocol with AWS's own definition that suits the scale up requirements. Trainium 2 uses NeuronLinkv3, which corresponds to PCIe 5, and Trainium 3 uses NeuronLinkv4, which corresponds to PCIe 6.

For Trainium 4, NeuronLinkv5 will be a bit more complicated as AWS aims to support both UALink and NVLink protocols. Off the package, Trainium 4 will offer UALink protocol running at 128G based on PCIe 7. This will be the basis of NeuronLinkv5 and all the scale up connections coming out of the Trainium 4 packages. For NVLink Fusion integration, the UALink protocol signal will be converted into NVLink lanes via an external gearbox off the Trainium 4 package. Below we will discuss the content each supplier plays in the NeuronLinkv5 UALink or NVLink protocol as well as how NPO ties into all of this.

Therefore, there will be two versions of the Trainium4 rack:

- UALink running at 128G based on PCIe Gen7
- PCIe Gen7 gearboxed into NVLink Fusion running at 200G per lane

The NPO modules are placed on the PCB via "[Si-Fly](#) CPX" connectors from Samtec close to the Trainium 4 package. The "[Si-Fly](#) CPX" connector socket is speced up to 224Gbs coconnectivity and 128DPs each. The connectors can support both fly over cables and NPO modules, which means AWS also have the flexibility to decide on NPO later given the flexibility. In the image below, Samtec demo their CPX connector under the CPO or CPC form factor, yet, the CPX connector also supports NPO and NPC (flyover cables) form factor. The use of SiFly will follow Amazon's design, which is the UALink rack, while Nvidia will likely use another solution.

samtec**SI-FLY HD****PAM4
224
Gbps**

Source: Samtec

Trainium4 Scale-up Network Suppliers

The ultimate production mix among the two rack types is highly uncertain, and one track could even get cut, but for now we pencil in a 50%/50% UALink/NVLink split over the product lifetime.

For UALink – we think Astera Lab will be the primary vendor, bundling its I/O chiplet, its UALink switch as well as an NPO solution. Astera has been pulling its timelines for NPO/CPO – late last year it had penciled in 2028-29 for scale-up optics, while by June 2026 it was expecting first optical revenue by 2027 – with NPO revenue in 2H27 for rack-to-rack connectivity, followed by CPO scale-up by 2028 and after.

Similarly, for the NVLink-based SKU, Astera will be providing the I/O chiplet with a gearbox, while Marvell will tackle the NPO module as well as the NVLink switch. In fact, a key revision to NVLink fusion in the past few months means that Nvidia is no longer the only place to get an NVSwitch, creating an opening for Marvell. Indeed, on the 4Q F1/26 earnings call, “NPO” entered Marvell’s vocabulary for the first time, and the company introduced a broader “scale-up optics” category guided to roughly \$300M for F1/28 – more than double the prior outlook of \$150M which was solely based on Celestial AI. It is important to highlight that the design of the rack will be

heavily driven by Nvidia to ensure successful implementation of the NVLink scale-up network, and as such, Nvidia may have the final word on whether NPO is used at all. For now, we pencil in the use of NPO for the NVLink SKU.

Though the \$500M quarterly revenue run rate for Celestial exiting F1/28 guided during the acquisition announcement was linked to a Tier 1 customer win and for a scale-up product (thus aligning with Trainium 4), subsequent commentary has highlighted CelestialAI's EAM based technology as one solution of a broad portfolio of NPO and CPO implementations that also include MRMs and MZMs and we think that Celestial's EAMs will not be used in the Trainium4 NPO solution. Instead, we expect to see Celestial products could be deployed in applications that utilize the photonic bridge or could include deployments of the Photonic Fabric Memory Appliance (PFMA) given that Marvell is well into the design phase for various hyperscaler CXL / memory attach products.

Trainium4 Scale-up Network Architecture

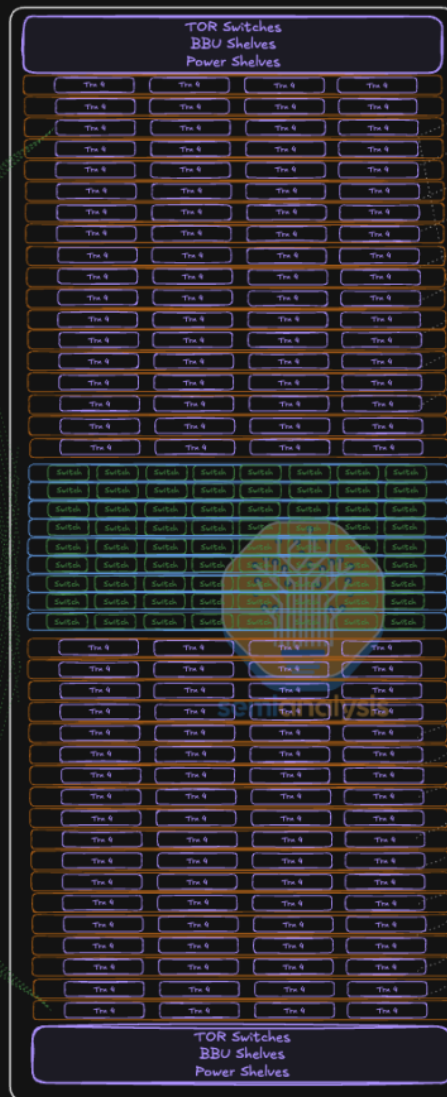
We think that Trainium4 will mirror the rack design of the Trainium3 NL72×2 (Teton3 MAX), only with greater density at 36 compute trays, with 144 XPUs per rack vs the 72 GPUs per rack of the Teton3 MAX. It will use backplane to connect to switch trays, could also use PCB traces to connect to smaller switch chips on the compute tray, and finally will use NPO to connect XPUs to distant switch trays in the other racks.

Trainium4 NL144x8 Switched

Rack 1

Rack 2

Each GPU connects via Backplane to 36 switch chips across 9 Switch trays.



GPUs use NPO to connect to 7 other racks.

○ ○ ○

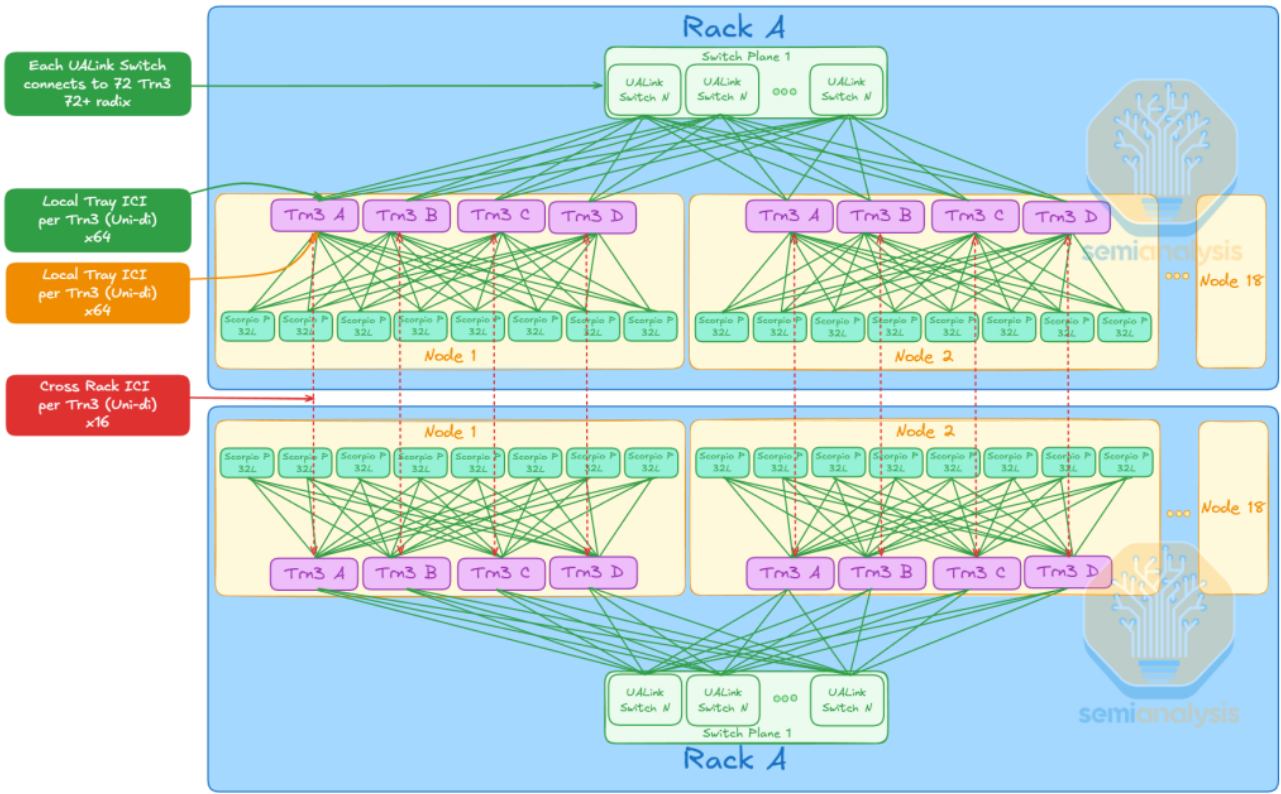
Rack 8



Source: [SemiAnalysis AI Networking Model](#)

We are still finalizing our view on what the scale-up networking topology could look like, but the [scale-up network topology employed in the Trainium3](#) could be a starting point. We should expect that Trainium4's topology could include smaller switches on the compute tray itself in addition to a backplane that connects to a set of larger bandwidth switches. Again, the main evolution would be that each Trainium4 is now also connecting via NPO to distant scale-up switches in other racks.

Trainium3 NL72x2 Switched NeuronLinkv4 Scale Up Topology
Scorpio X (72+ Port UALink)



Source: SemiAnalysis AI Networking Model

While we are still finalizing our view on what the scale-up topology would look like, we see 144 GPUs delivering 4.72 Pbit/s uni-di of total scale-up bandwidth in the UALink configuration and 3.69 Pbit/s in the NVLink configuration, which equates to 6 and 4.75 of 3.2T-equivalent NPO OEs per GPU.

Optical Engine and ELS Content per Rack Estimates

	Unit	VRU NVL576 CPO	Trainium 4 NVLink ¹	Trainium 4 UALink ¹
Logical GPUs per Compute Tray	GPUs	4	4	4
Compute Blades per Rack	Blades	18	36	36
Logical GPUs per Rack	GPUs	72	144	144
L1 Scale-up Switch Chips per Switch Blade	Switches	6	8	8
Switch Blades per Rack	Blades	9	9	9
L1 Scale-up Switch Chips per Rack	Switches	54	72	72
L2 Scale-up Switches per Scale-up Domain	Switches		0.0	0.0
L2 Scale-up Switches allocated per Rack	Switches		0.0	0.0
OE Attached to GPU (3.2T Equiv.)	OEs	0.0	4.8	6.0
OE per L1 Switch Chip (3.2T Equiv.)	OEs	3.0	9.5	12.0
OE per L2 Switch Chip (3.2T Equiv.)	OEs			
ELS Attach to GPU (6.4T Equiv.)	ELSFPs	0.0	2.38	3.00
ELS per L1 Switch Chip (6.4T Equiv.)	ELSFPs	1.50	4.8	6.0
ELS per L2 Switch Chip (6.4T Equiv.)	ELSFPs			
GPU OE per Rack (3.2T Equiv)	OE/Rack	0	684	864
L1 Switch OE per Rack (3.2T Equiv)	OE/Rack	162	684	864
L2 Switch Allocated OE per Rack (3.2T Equiv)	OE/Rack	0	0	0
External Laser Source per Rack (6.4T Equiv.)	ELSFP/Rack	81	684	864
Optical Engines per GPU (3.2T Equiv.)	OE/GPU	2.25	9.5	12.0
External Laser Source per GPU (6.4T Equiv.)	ELSFP/GPU	1.125	4.8	6.0

1. Specifications are all preliminary - no architectures have yet been confirmed.

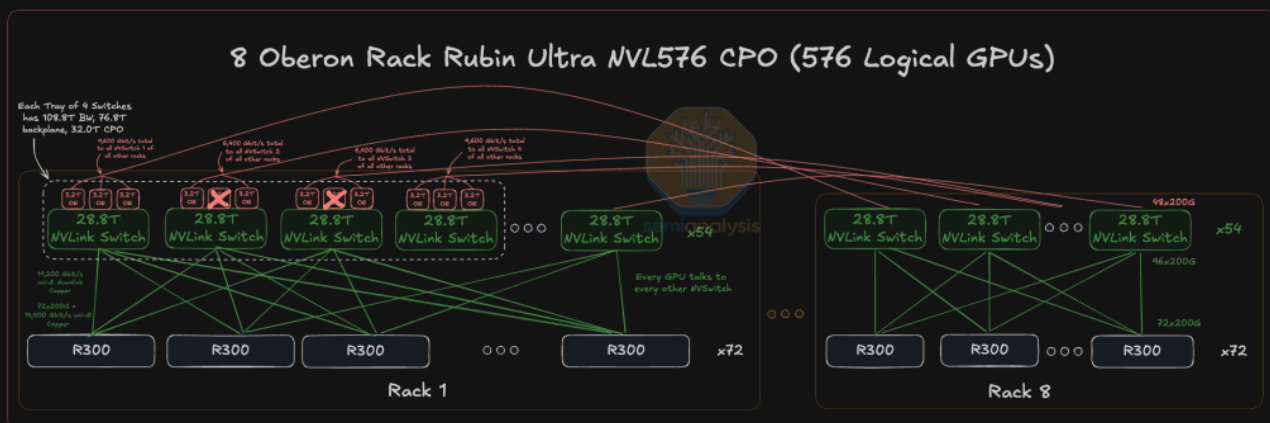
Source: [SemiAnalysis AI Networking Model](#)

Nvidia Scale-up NPO

We have discussed the various yield issues and fiber attach challenges Nvidia has been facing on its CPO program – these also impact the Vera Rubin NVL576 8-rack scale-up CPO project and may even carry over to a future OCI-MSA style DWDM-based CPO that would have otherwise made its debut in Feynman. This is why Nvidia has been contingency planning by ramping a potential NPO implementation that focuses on 200G or 400G PAM4 also using DWDM.

Nvidia's VRU NVL72 CPO architecture will employ the QM5 NVLink Switch with 3× 3.2T Optical engines as a building block. Each switch will have 19.2Tbit/s uni-di of connectivity using a copper backplane connecting to all

GPUs within the same rack and will employ either two or three optical engines to connect with at most 9.6Tbit/s uni-di via CPO to NVLink switches in other racks. Each compute tray of 4 NVSwitch chips will have a total of 108.8Tbit/s uni-di of bandwidth, 76.8Tbit/s connecting within the rack over the backplane, and 32.0Tbit/s via a total of 10 actively used 3.2T Optical Engines to connect to the other 7 racks in the system. Only 10 out of the 12 optical engines available in each switch tray are used, giving rise to the irregular CPO bandwidth coming off each NVLink switch – two connect with 6.4Tbit/s of bandwidth and two with 9.6Tbit/s.



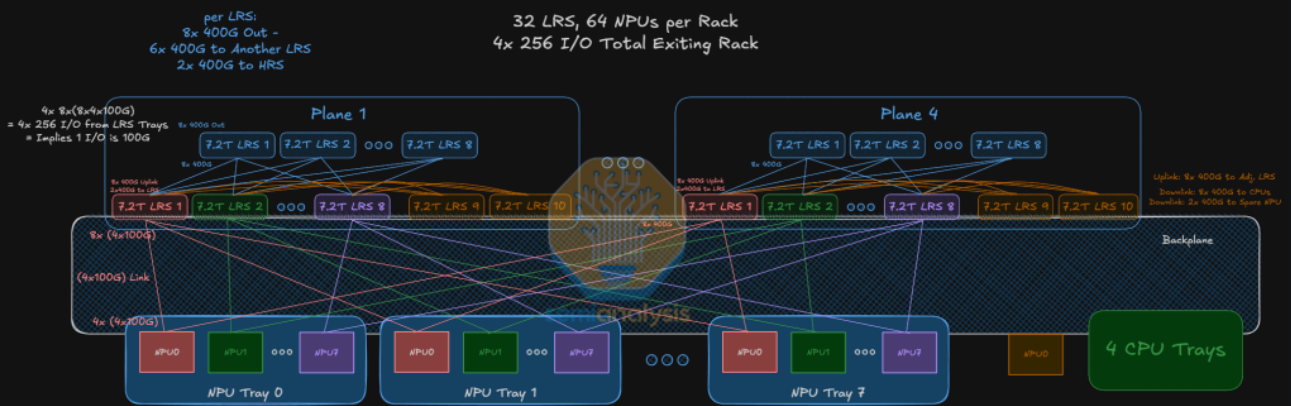
Source: [SemiAnalysis AI Networking Model](#)

If the above NPO implementation is ready in time for the VRU NVL576, then Nvidia could look to substitute NPO optical engines in place of the CPO optical engines. It is also quite possible given the high production risk that Nvidia shelves the cross-rack scale-up design for Vera Rubin Ultra and continues with Oberon style racks until Feynman.

Huawei NPO

Huawei's Atlas 950 SuperPod implements a multi-stage dragonfly with a scale-up world size of ~8k NPUs: 8× 1,024-NPU UB-Mesh-Pods connected by high-radix switches (HRS), each UB-Mesh-Pod comprising 16 racks connected in a 4D FullMesh, and each rack pairing 64 NPUs with 7.2T lower-radix switches (LRS) across 4 planes. We covered the architecture in depth [here](#), including our estimate of \$192 per NPU of in-rack copper backplane content.

Rack Architecture



Source: SemiAnalysis AI Networking Model

Unlike Trainium4, where NPO reaches all the way to the GPU, Huawei keeps the NPU side copper – of each NPU's 18x 400G ports (7.2 Tbit/s scale-up bandwidth), seven connect to other NPUs on the same tray over flyover cables, seven connect across trays in the intra-rack 2D mesh over the backplane, and four connect up to the LRS layer over the backplane as well. NPO comes exclusively at the LRS-to-LRS layer where every switch-to-switch link within and exiting the rack is optical.

Intra-rack 2D Mesh Topology

NPU Rack, 64 NPUs per Rack

Source: [SemiAnalysis AI Networking Model](#)

Ascend 950 In-Rack Copper Content			
Item	Ascend 950		
	Unit Cost	Quantity per Server	Extended Price
GPUs per Server	N/A	8,192	N/A
<i>112Gb/s 48DP Connector (Female)</i>	\$22	8,192	180,224
<i>112Gb/s 48DP Connector (Male)</i>	\$66	8,192	540,672
<i>Single DP Copper Cable</i>	\$1	393,216	393,216
Total Backplane Content			\$1,114,112
<i>Mid-Board Connectors</i>	N/A	0	0
<i>Flyover Cables</i>	\$8	57,344	458,752
Total Compute Tray Content			\$458,752
Total Backplane Content per XPU			\$1,572,864
			\$192

Source: [SemiAnalysis AI Networking Model](#)

The LRS network is a two-level hierarchy per plane – 10 L1 switches and 8 L2 switches in each plane which comes to 72 LRS switches per rack. The first eight L1 switches (LRS 1-8) are NPU-facing leaves – each takes 8× 400G downlinks from NPUs through copper and drives 10× 400G of uplinks through NPO – eight uplinks in a full mesh to the plane's L2 switches and one link each to the intra-plane LRS switches (LRS 9 and LRS 10). Each L2 switch takes 8× 400G downlinks from the L1 leaves and drives 8× 400G up to the L2 switches in the other 15 racks of the UB-Mesh-Pod, all through NPO.

Calculating NPO attach rate, on the L1 side, 256 leaf uplinks and 128 hub links give 384 ports, on the L2 side, 256 downlinks and 256 rack exits give 512 ports – which in total gives us 896× 400G optical port-ends per rack, or 358.4 Tbit/s of NPO bandwidth. At 3.2T per optical engine, that is 112 optical engines per rack – an attach ratio of 1.75× 3.2T-equivalent OEs per NPU.

Huawei NPO Content per Rack		
	Units	Quantity
LRS 1-8 Uplinks to L2	#	256
LRS 1-8 Downlinks from LRS 9-10	#	64
LRS 9-10 Uplinks to LRS 1-8	#	64
Total L1 Links	#	384
LRS 1-8 Downlinks from LRS 1-8	#	256
LRS 1-8 Uplinks to Other Racks	#	256
Total L2 Links	#	512
Total Links	#	896
BW per Link	Gbit/s	400
NPO BW per Rack	Tbit/s	358.4
NPO BW per Module	Tbit/s	3.2
3.2T-equivalent NPO per Rack	#	112
3.2T-equivalent NPO per GPU	#	1.75

Source: [SemiAnalysis AI Networking Model](#)

Scale-out NPO Projects

Meta NPO

Meta has spent considerable time and energy highlighting the advantages of Broadcom's CPO solution. Their ECOC paper published reliability data across 1,409,000 400G-port device hours of operation for Bailly 51.2T CPO Switches. Its talk at ECOC presenting the paper went even further, presenting results for up to 15M 400G-port device hours, showing no uncorrectable codewords (UCWs) for the first 4M 400G-port device hours as well as a 0.5-0.1M device hour mean time before failure (MTBR) for 400G 2xFR4 transceivers (550k for 2xFR4 globally) vs 2.6M device hour MTBF for CPO.

A follow-up update at OFC 2026 highlighted that testing had now reached 40M 400G-port device hours, and pointing out the >20M 400G-port device hours MTBF.

CPO Evaluation Infrastructure - Updated March' 26

Optics	Operating Temperature	Device Hours (400G)	MTBF
2x400G FR4 pluggable	40°C	~8M	0.71M
CPO Phase 1	40°C	>40M	1.47M
CPO Phase 1 (excluding ELSFP issue*)	40°C	>40M	8.2M (>10X improvement)
CPO Phase 1 (non-serviceable)	40°C	>40M	>20M

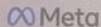
* ELSFP issue isolated to laser driver circuit SMT component

Data shows that CPO ports are more reliable than pluggable module ports, potentially due to:

- Reduced number of interfaces and deeper level of integration
- Flexibility in co-design of the system with improved operating condition for key components (ie lasers)
- Integration of optics within the platform allows system level test and screening in the manufacturing line
- Reduced human intervention

CPO Phase 2 testing revealed too few failures at 50M hrs to report a confident MTBF value

Field serviceable ELSFP modules improve deployed reliability and decrease blast radius concern



Detailed reporting in Session M4B | S1: Datacom Subsystems and Systems "Optical Engines": Monday

Source: Meta

There has been some follow thru on CPO switches – Celestica will be manufacturing 102.4T switches for Meta, with production ramping in 2H27. However, Meta is also simultaneously pursuing an NPO switch program based on Broadcom's 102.4T Tomahawk 6 Switch. In theory, CPO should be able to provide a better channel than NPO and thus eventually superior power and BoM cost – but NPO's advantage lies in optical engine module vendor flexibility and the cost savings from unbundling the optical engines from the switch ASIC vendor and effectively paying a 30-40% gross margin instead of a 70% + margin. Another factor in favor of deploying NPO-based solutions at this stage is the ongoing yield and production issues for CPO optical engines built using TSMC's COUPE.

Unlike the Davisson CPO platform, this NPO approach allows a multivendor sourcing approach for optical engines, and Meta will almost certainly be working with its current transceiver vendors such as Innolight, Eoptolink, Coherent, and others. This NPO switch also allows operators to mix and match various types of optical engines. For instance, some hyperscalers are

opting to deploy FR optics in some leaf to spine links to simplify fiber deployments and reduce fiber costs – so a switch that could utilize DR NPO for downlink and FR NPO for uplink could be useful.

Meta is taking its CPO program planning one step further and has been exploring CPO Optical Engines and ELSFPs that support OCI MSA Gen1 specs – with 8 different wavelengths per 12.8T ELSFP. The light source can be split across two 6.4T Optical Engines, with each 6.4T OE providing 32 lanes of 200G each, and each lane in turn built up with 4 lambdas modulated at 50G NRZ each. The use of OCI MSA means it will almost certainly be for a scale-up role, possibly with the Tomahawk Ultra 2 102.4T switch.

In pursuing an OCI MSA based CPO switch, Meta may be aiming to bootstrap the OCI MSA switch and OE ecosystem and could be looking to see how much traction it can get and whether it can not only build its own MTIA scale-up networks around the OCI MSA but also incorporate this into merchant silicon systems like that of AMD's, which is also an OCI MSA consortium member. If there is good buy in from vendors and suppliers, then this could potentially give rise to broader adoption of OCI MSA beyond Meta, but time will tell how much traction they get. If Meta has enough clout, they might even convince Broadcom to work with them to create a pluggable CPO version, which would doubtless be welcomed by hyperscalers and ASIC vendors.

Meta CPO and NPO Content TAM (Install Basis Convention)						
	Units	2026	2027	2028	2029	2030
Meta CPO - DR Optics	102.4T Switches	0	6,500	24,500	11,500	0
Meta NPO	102.4T Switches	0	15,147	41,918	42,112	11,457
Meta CPO - OCI MSA	102.4T Switches	0	0	35,000	45,000	5,000
Total 102.4T Equiv Switches	102.4T Switches	0	21,647	101,418	98,612	16,457
Meta CPO - DR Optics	OEs	0	208,000	784,000	368,000	0
Meta NPO	OEs	0	484,704	1,341,360	1,347,581	366,638
Meta CPO - OCI MSA	OEs	0	0	1,120,000	1,440,000	160,000
Total 3.2T Equiv OEs	OEs	0	692,704	3,245,360	3,155,581	526,638
Meta CPO - DR Optics	ELsSs	0	104,000	392,000	184,000	0
Meta NPO	ELsSs	0	242,352	670,680	673,790	183,319
Meta CPO - OCI MSA	ELsSs	0	0	560,000	720,000	80,000
Total 6.4T Equiv ELSFP	ELsSs	0	346,352	1,062,680	857,790	183,319

Source: [SemiAnalysis AI Networking Model](#)

Arista's Interconnect Options – XPO vs. NPO

We [discussed the merits of the XPO MSA](#) as an answer to CPO right after the MSA dropped in March 2026. The large form-factor pluggable promotes better thermal management by enabling LPO and could quadruple the bandwidth density of a faceplate.

The XPO solution retains the pluggable model, but in a new, much larger form factor, packing 64×200G lanes in the same module, all while allowing better thermal management with liquid cooling. This would only require 16×12.8T transceivers per switch; Arista claims that they can deliver a total of 204.8T on a 10U front panel and keeping the benefits of the pluggable: being hot swappable.

XPO aims to come to market alongside, Tomahawk 7, which will double the port radix to 1,024 200G SerDes compared to Tomahawk 6, which has 512 200G SerDes. Switch vendors that are building Tomahawk 7 based switches, but are stopping short of non-pluggable CPO, therefore have three available options: Use XPO, NPO/Pluggable CPO and OSFP transceiver cages.

Arista Networks is clearly focused on bringing XPO to market for Tomahawk 7 switches, but it would not be surprising if they are also entertaining NPO-based solutions also. Both achieve the objective of a much denser faceplate, and both open the door to eliminating DSPs. The main difference boils down to where it wants to take implementation risk, in getting optics to work or in getting a longer electrical channel to work with Linear optics.

An NPO solution would allow for a very high faceplate density as the optical engine is now placed on the PCB, only leaving fiber connectors on the front of the switch. However, servicing a failed optical engine requires opening the box and reseating or replacing a socketed engine – far more intrusive than hot-swapping a front-panel pluggable, though external laser sources remain field-replaceable.

The OSFP transceiver-based option will still likely remain popular among operators that have an established, low-cost 1.6T transceivers supply chain, despite the large switch form factors that this implies.

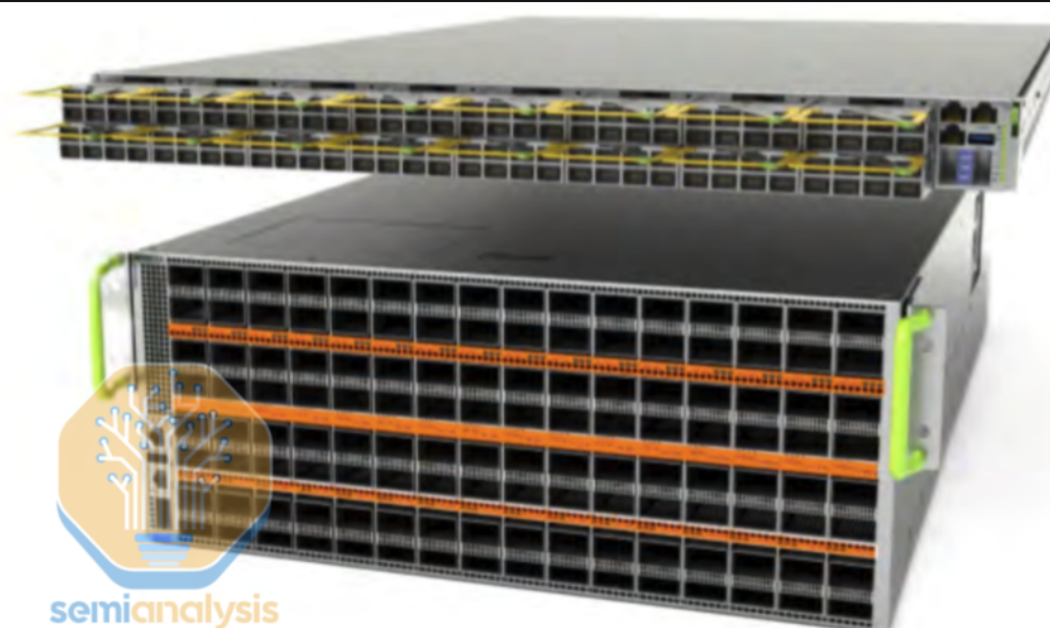


Figure 3: Illustration of 204.8Tbps switch with XPO (top) and OSFP (bottom) showing the 4X density improvement

Source: [SemiAnalysis AI Networking Model](#), Arista Networks

Google NPO

Google has in the past been very pluggable optics focused, but some have claimed that they are also exploring the use of NPO in future TPU versions, with considerable orders already placed for delivery in late 2026 and early 2027! We think Google could explore NPO later for its Virgo network, but we don't see a role for NPO in the next few years and think these reports probably stem from some kind of misunderstanding.

The argument here is that NPO could be used to enable higher density interconnect on the compute trays, replacing the use of ACCs and DACs. However – the TPUv8i (Boardfly) only requires 2× 1.6T ACCs, 1× 1.6T DAC and 1× 1.6T pluggable optical transceiver, employing another 1.6T over flyover cables and 1.6T over PCB traces to achieve 9.6T of bandwidth. Meanwhile, the TPU V8t (3D Torus) also has 9.6T of scale-up bandwidth per TPU, carried over PCB as well as 4 OSFP cages per TPU to which ACCs or DACs or Pluggable transceivers can be inserted depending on where in the

3D Torus topology that TPU is. Again, this is hardly pushing the limits on faceplate density – so the use of NPO here seems to only add risk and complexity, with the main positive, NPO's greater range, not needed for the copper-based links.

Some also say that Google could use NPO for Virgo. The first instances of the Virgo network are currently being deployed and will be in use with the TPUv8t, which means that Virgo network is not using XPO or NPO as of now, and current scale-out bandwidth needs can easily be covered by current pluggable solutions. Google is likely to keep the Virgo network for scale-out with the upcoming version of the TPUs, mainly with Humufish (TPUv9t) and Icefish (TPUv10t). Current scale-out bandwidth for Zebrafish generation (TPUv8t) sits at 800G per TPU but is very likely to increase over time with next generations, which could require the Virgo network to adopt XPO with Tomahawk 6 and TPUv9t, but NPO would only come in with Tomahawk 7 and TPUv10t in the coming years.

Here, the benefits are like what has been mentioned above in our XPO vs NPO discussion. NPO is another way to implement linear optics, but the tradeoff vs XPO is taking on different risks – XPO uses CPC to deliver a cleaner electrical channel to the faceplate, but the attenuation due to distance is unavoidable. NPO cuts that electrical channel distance considerably, but it introduces the risk of ramping optical production and reliability risk.

CPO and NPO Market Estimates Updated

We have updated our CPO and NPO market projections to incorporate the developments outlined.

CPO and NPO TAM Build-Up (Install Base Convention)						
	Units	2026	2027	2028	2029	2030
NPO Scale-Out 3.2T Equiv OEs	Units	0	484,704	1,341,360	1,347,581	366,638
NPO Scale-Up 3.2T Equiv OEs	Units	399,000	1,312,500	13,430,552	40,585,385	24,841,385
NPO 3.2T Equiv OEs	Units	399,000	1,797,204	14,771,912	41,932,966	25,208,024
CPO Scale-Out 3.2T Equiv OEs	Units	11,847	1,407,865	4,776,733	6,362,874	9,081,450
CPO Scale-Up 3.2T Equiv OEs	Units	0	0	2,731,308	15,182,183	46,810,600
CPO 3.2T Equiv OEs	Units	11,847	1,407,865	7,508,040	21,545,057	55,892,050
Total 3.2T Equivalent Optical Engines - CPO + NPO	Units	410,847	3,205,069	22,279,953	63,478,023	81,100,074

Source: [SemiAnalysis AI Networking Model](#)

CPO and NPO TAM Build-Up (Install Base Convention)						
	Units	2026	2027	2028	2029	2030
NPO Scale-Out 3.2T Equiv OEs	Units	0	484,704	1,341,360	1,347,581	366,638
NPO Scale-Up 3.2T Equiv OEs	Units	399,000	1,312,500	13,430,552	40,585,385	24,841,385
NPO 3.2T Equiv OEs	Units	399,000	1,797,204	14,771,912	41,932,966	25,208,024
CPO Scale-Out 3.2T Equiv OEs	Units	11,847	1,407,865	4,776,733	6,362,874	9,081,450
CPO Scale-Up 3.2T Equiv OEs	Units	0	0	2,731,308	15,182,183	46,810,600
CPO 3.2T Equiv OEs	Units	11,847	1,407,865	7,508,040	21,545,057	55,892,050
Total 3.2T Equivalent Optical Engines - CPO + NPO	Units	410,847	3,205,069	22,279,953	63,478,023	81,100,074
Scale-out FAUs and Assembly - CPO + NPO	Units	11,847	1,892,569	6,118,093	7,710,455	9,448,088
Scale-up FAUs and Assembly - CPO + NPO	Units	399,000	1,312,500	16,161,860	55,767,568	71,651,985
Fiber Attach Unit (FAUs) and Assembly	Units	410,847	3,205,069	22,279,953	63,478,023	81,100,074
Scale-out 6.4T ELS - CPO + NPO	Units	5,924	946,284	3,619,046	4,575,227	4,804,044
Scale-up 6.4T ELS - CPO + NPO	Units	199,500	656,250	8,080,930	27,883,784	35,825,993
6.4T External Laser Source (ELS)	Units	205,424	1,602,534	11,699,976	32,459,011	40,630,037
Scale-out Shuffle Box	Units	304	33,005	133,414	179,032	253,788
Scale-out Optical Engines	\$M	\$1	\$189	\$612	\$771	\$945
Scale-up Optical Engines	\$M	\$40	\$131	\$1,616	\$5,577	\$7,165
Total Optical Engines	\$M	\$41	\$321	\$2,228	\$6,348	\$8,110
Scale-out FAUs	\$M	\$2	\$331	\$1,071	\$1,349	\$1,653
Scale-up FAUs	\$M	\$32	\$115	\$1,159	\$5,438	\$6,912
Fiber Attach Unit (FAUs)	\$M	\$34	\$446	\$2,230	\$6,787	\$8,565
Scale-out 6.4T ELS	\$M	\$2	\$331	\$1,071	\$1,349	\$1,653
Scale-up 6.4T ELS	\$M	\$70	\$230	\$2,828	\$9,759	\$12,539
6.4T External Laser Source	\$M	\$72	\$561	\$3,899	\$11,109	\$14,193
Scale-out Shuffle Box	\$M	\$1	\$83	\$334	\$448	\$634
Total Key Component TAM	\$M	\$148	\$1,410	\$8,691	\$24,691	\$31,502
.. Of which: Scale-out	\$M	\$6	\$934	\$3,087	\$3,917	\$4,886
.. Of which: Scale-up	\$M	\$142	\$475	\$5,604	\$20,774	\$26,616

1. GPU Shipment figures are preliminary and will likely change once the AI Accelerator Model initiates demand forecasts through 2030.

Source: [SemiAnalysis AI Networking Model](#)

Disclaimers

Analyst Certifications and Independence of Research.

Each of the analysts whose names appear in this report hereby certify that all the views expressed in this Report accurately reflect our personal views about any and all of the subject securities or issuers and that no part of our compensation was, is, or will be, directly or indirectly, related to the specific recommendations or views in this Report. SemiAnalysis LLC (the "Company") is an independent equity research provider. The Company is not a member of

the FINRA or the SIPC and is not a registered broker dealer or investment adviser. SemiAnalysis has no other regulated or unregulated business activities which conflict with its provision of independent research.

Limitation Of Research And Information.

This Report has been prepared for distribution to only qualified institutional or professional clients of SemiAnalysis LLC. The contents of this Report represent the views, opinions, and analyses of its authors. The information contained herein does not constitute financial, legal, tax or any other advice. All third-party data presented herein were obtained from publicly available sources which are believed to be reliable; however, the Company makes no warranty, express or implied, concerning the accuracy or completeness of such information. In no event shall the Company be responsible or liable for the correctness of, or update to, any such material or for any damage or lost opportunities resulting from use of this data. Nothing contained in this Report or any distribution by the Company should be construed as any offer to sell, or any solicitation of an offer to buy, any security or investment. Any research or other material received should not be construed as individualized investment advice. Investment decisions should be made as part of an overall portfolio strategy and you should consult with a professional financial advisor, legal and tax advisor prior to making any investment decision. SemiAnalysis LLC shall not be liable for any direct or indirect, incidental or consequential loss or damage (including loss of profits, revenue or goodwill) arising from any investment decisions based on information or research obtained from SemiAnalysis LLC.

Reproduction and Distribution Strictly Prohibited.

No user of this Report may reproduce, modify, copy, distribute, sell, resell, transmit, transfer, license, assign or publish the Report itself or any information contained therein. Notwithstanding the foregoing, clients with access to working models are permitted to alter or modify the information contained therein, provided that it is solely for such client's own use. This Report is not intended to be available or distributed for any purpose that would be deemed unlawful or otherwise prohibited by any local, state, national or international laws or regulations or would otherwise subject the Company to registration or regulation of any kind within such jurisdiction.

Copyrights, Trademarks, Intellectual Property.

SemiAnalysis LLC, and any logos or marks included in this Report are proprietary materials. The use of such terms and logos and marks without the express written consent of SemiAnalysis LLC is strictly prohibited. The copyright in the pages or in the screens of the Report, and in the information and material therein, is proprietary material owned by SemiAnalysis LLC unless otherwise indicated. The unauthorized use of any material on this Report may violate numerous statutes, regulations and laws, including, but not limited to, copyright, trademark, trade secret or patent laws.

Private Investment Disclosure.

Each analyst and/or author contributing to this Report hereby discloses that they may hold, directly or indirectly, private investments in companies, sectors, or asset classes discussed or referenced herein. Such holdings, if any, are disclosed externally with a third-party vendor. SemiAnalysis LLC requires that no author or analyst allow their personal investment positions to influence the views, opinions, analyses, or conclusions expressed in this Report. Readers should be aware that the existence of such holdings, even where disclosed, may present a potential conflict of interest, and are encouraged to weigh this accordingly when evaluating the contents of this Report.

ETF and Regulatory Status Disclaimer.

From time to time, The Company may collaborate with asset managers, issuers, sponsors, index providers or other financial institutions in connection with the development, launch or support of exchange-traded funds ("ETFs") and other investment products, including by providing research, market analysis, index-related research or other research related services. The Company does not sponsor, endorse, sell or promote any ETF, and makes no representation or warranty, express or implied, regarding the advisability of investing in any such ETF. In connection with the services described herein, the Company does not act as a fiduciary to any investor, fund, sponsor, issuer or other party. The Company is not registered with the U.S. Securities and Exchange Commission or any state securities authority as an investment adviser or exempt reporting adviser. The Company's role is limited to providing research and related services, and it does not provide

individualized investment advice, portfolio management, brokerage services or other services requiring registration as an investment adviser.