

From Electrons to Light

- Copper is running out of road inside the AI data center — light is becoming the default plumbing, not a niche upgrade.
- Three zones, three timelines: Scale-Across converted first, Scale-Out is where the market sits today, Scale-Up is next to grow.
- Our playbook: own the pluggable-era winners now, position early for the names CPO will reward later.

A Dawn of Optical Connectivity

AI investment has moved through distinct waves, each defined by a different bottleneck — first compute, then memory, and now what happens between thousands of GPUs once they have to behave as one machine. That problem splits into power and connectivity, and connectivity is the subject of this paper. **Our thesis is simple: copper, the wire that has carried digital signals for over a century, is running out of road inside the AI data center.** Photonics — light instead of electricity — is moving from a niche technology into the default plumbing of every new AI cluster.

Physics of the Copper wall

Push enough speed through copper and it turns against itself: the signal degrades and runs out of reach, while the wasted power shows up as heat you then have to pay to remove. Neither problem is a design flaw to engineer around — it's physics that compounds with every new GPU generation, hitting every connection inside the data center eventually, just not all at once. That's the whole reason photonics, transferring the data via light, is a structural requirement here, not just an upgrade.

Three network zones, and the slow march to CPO

Not every connection inside an AI data center faces the same problem at the same time. There are three types of connection: 1) **Scale-Across**, connecting data centers to each other, which has run on light for decades. 2) **Scale-Out**, rack to rack, made that same jump from copper years ago and is where today's market volume actually sits — running on pluggable modules at 800G moving to 1.6T, with the industry already working through cost and power problems as it moves toward tighter chip-level integration with CPO. 3) **Scale-Up**, GPU to GPU within a rack, which is still mostly copper — but is set to become the biggest growth story of the three, as pod sizes outgrow a single rack and AI inference itself starts demanding GPUs talk to each other more often, not less.

Who benefits: A Layer-by-Layer map

In this paper, we separate the photonics supply chain into seven layers — from the raw InP wafer every laser starts with, through the electrical ICs and photonic foundry, to module assembly and the system-level decision of whether a link stays copper or goes optical. It's a genuinely fragmented chain, unlike the way the GPU market clusters around NVIDIA, with each layer carrying its own materials, factories, and set of key beneficiary companies.

Investment Implications

Our playbook: own the pluggable-era winners now, position early for the names CPO will reward later. Near-term, that's the assembly leaders and the chips still stretching copper's last few meters; medium-to-longer-term, it's the foundries and packaging specialists betting CPO becomes a semiconductor business, not just an assembly one. Across both legs — needed no matter which way the architecture goes — sit the substrate suppliers and the vertically-integrated laser makers. Track one tripwire above all: how fast CPO moves from qualification to real shipping volume.

Analyst

Suwat Sinsadok, CFA, FRM, ERP
suwat.s@globlex.co.th,
+662 687 7026

Assistant Analyst

Natthapol Changsamlee

From Electrons to Light

A Dawn of Optical Connectivity

AI investment has moved through distinct waves, each defined by a different bottleneck. The first was compute — the race for more and faster GPUs (Graphics Processing Units, the chips that do the parallel math behind AI), which made NVIDIA the world's most valuable company. The second was memory — HBM (High Bandwidth Memory), the fast on-package memory that lets an AI model hold a longer conversation without forgetting the start of it.

With each bottleneck solved, the next one comes into view. Today it is what happens between chips, once thousands of GPUs have to behave as one machine. That problem splits into two parallel fights: power — delivering and cooling enough electricity to increasingly dense racks — and connectivity, the wiring itself, which is the subject of this paper.

Our thesis is simple: **copper, the wire that has carried digital signals for over a century, is running out of road inside the AI data center.** At today's speeds, copper loses too much signal, wastes too much power as heat, and cannot reach far enough. Photonics — using light instead of electricity to carry data — is moving from a niche technology into the default plumbing of every new AI cluster.

To get the most out of this paper, we recommend getting familiar with these technical terms first.

Exhibit 1: Glossary Box

Term	Meaning
GPU	Graphics Processing Unit — the chip that does AI's parallel math.
ASIC	Application-Specific Integrated Circuit — a chip built for one job, e.g. a switch ASIC.
HBM	High Bandwidth Memory — fast memory stacked next to a GPU.
NVLink	NVIDIA's proprietary GPU-to-GPU links inside a rack.
ISI	Intersymbol Interference — signal pulses smearing into their neighbors.
InP	Indium Phosphide — the material that can actually emit light; silicon cannot.
SiPh	Silicon Photonics — using chip manufacturing to guide light, not generate it.
PIC	Photonic Integrated Circuit — a chip that guides and manipulates light.
Die	A single individual chip cut from a wafer, before it's packaged into a finished module.
DSP	Digital Signal Processor — cleans up a degraded signal.
TIA	Transimpedance Amplifier — The receive-side chip that turns a faint returning light signal into a readable voltage.
CW laser	Continuous Wave laser — an always-on beam, modulated separately.
EML	Electro-absorption Modulated Laser — the older design that combines the laser and its modulator on one InP chip
AEC	Active Electrical Cable — a copper cable with power-hungry signal-boosting chips built in, used to squeeze a bit more reach out of copper.
CPO	Co-Packaged Optics — the optical engine mounted on the same package as the switch chip.
NPO	Near-Packaged Optics — intermediate steps toward full CPO.
OSAT	Outsourced Semiconductor Assembly and Test — advanced chip packaging, harder than simple module assembly.
Fabless	A company that designs chips but doesn't own the factories that manufacture them — it pays a foundry to build what it designs.
Foundry	A factory that manufactures chips designed by other companies.

Sources: Globlex research

Physics of the Copper wall

Why copper worked, and why it stops working now

Think of an AI data center as one factory floor built to run a single enormous calculation. Compute (CPU, GPU, and ASIC) are the workforce. Memory is each worker's desk. Connectivity is the hallway that lets every worker talk to every other worker, instantly. A single GPU is trivial; the hard part is that a modern AI model runs on tens of thousands of GPUs, wired so tightly that the whole cluster has to act like one giant chip. That wiring is not a background detail — it is increasingly the product itself.

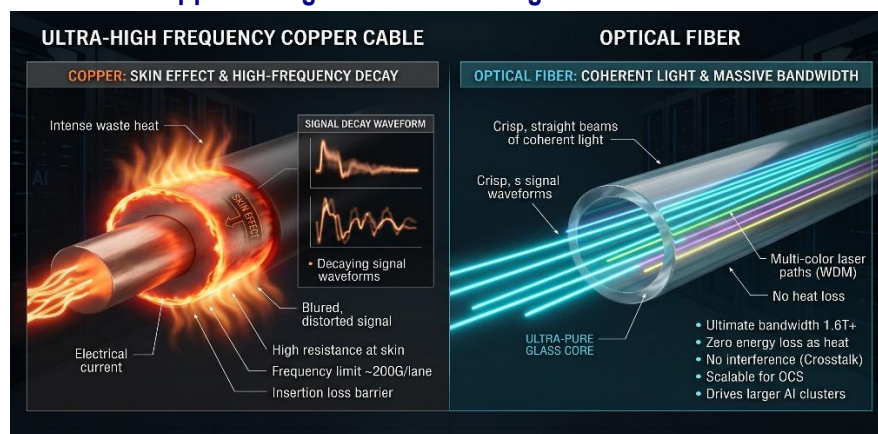
Copper is cheap and simple, and for most of computing history that was enough. Copper is not a bad conductor. The problem is what happens to any conductor once you push more speed through it — and AI racks now demand exactly that, because no GPU can sit idle waiting for its neighbor's answer. More speed means a higher-frequency signal, and that single change is the root of two separate problems, which are easy to conflate.

- **Problem A — the signal degrades and runs out of reach.** At high frequency, current crowds into the outer surface of the wire namely “**Skin effect**”, so the wire behaves as if it were thinner than it is, and resistance rises. The insulation around it also quietly absorbs more energy the higher the frequency climbs (dielectric loss). Both effects worsen with distance. Different frequencies inside the signal arrive slightly out of step, so pulses smear into their neighbors — this is ISI (Intersymbol Interference), a fair description of the signal literally colliding with itself.
- **Problem B — the wasted power becomes wasted heat.** This is a separate consequence of the same cause. As skin effect raises resistance, more energy is lost as heat inside the cable itself, rather than arriving as usable signal I^2R loss, where I is current and R is resistance). In a rack already fighting to remove heat, that waste is doubly expensive — once as electricity is paid for and unused, again as extra cooling load.

A shouted conversation down a narrowing, echoing tunnel is a fair picture of copper at high speed: stand further apart or talk faster, and your words blur together by the time they arrive while Light behaves more like a laser pointer than a shout — it does not blur or bleed energy into the tunnel wall the same way, so it can travel further and faster with far less waste.

That pressure is not static — it compounds every generation, and it has done so for a decade. It doesn't hit every connection inside the data center at once: some links are still short enough that copper holds today, others crossed this wall years ago. But directions are the same everywhere — this is why photonics is a structural requirement, not an upgrade.

Exhibit 2: Copper vs. Light: Where Each Signal Breaks Down



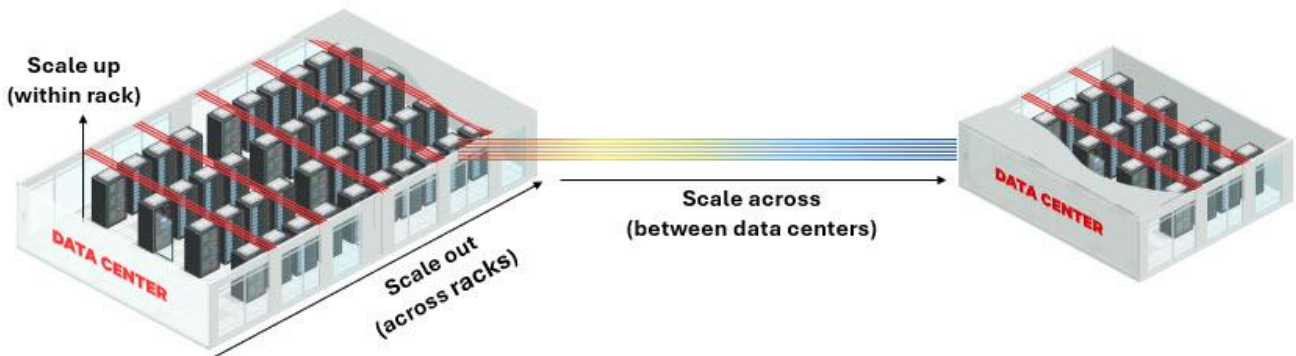
Sources: Globlex research

Three network zones, and the slow march to CPO

Three distances, three solutions

Not every connection inside an AI data center faces the same problem, so the industry does not solve them all at once. Think of it like an org chart: desk-to-desk, floor-to-floor, building-to-building.

Exhibit 3: The three network zones inside an AI Data Center



Sources: ciena.com

Exhibit 4: Scale-Up, Scale-Out, Scale-Across — Distance and Medium Compared

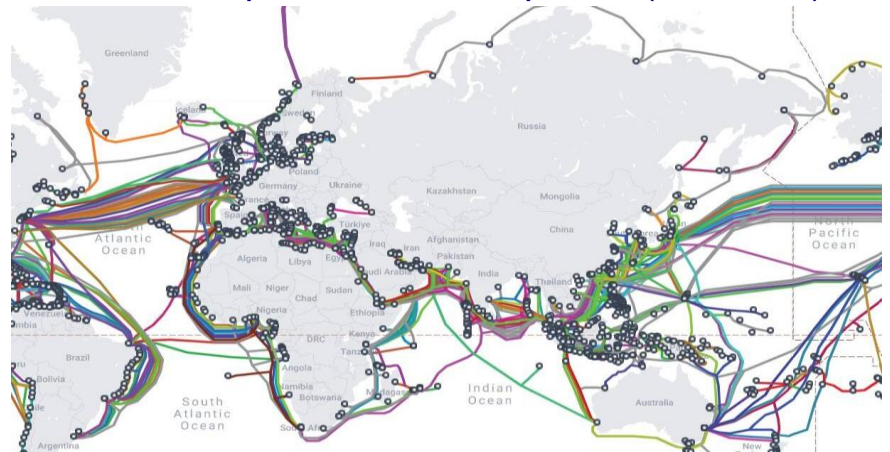
Zone	Connects	Distance	Medium today
Scale-Up	GPU to GPU, same rack (NVLink)	<1 meter	Still mostly copper
Scale-Out	Rack to rack	Meters to ~100m	Pluggable optics — 800G moving to 1.6T
Scale-Across	Data center to data center	Kilometers	Coherent optics — essentially 100% optical already

Sources: Globlex research

Scale-Across — never a copper story

Scale-Across connects one data center to another, across a metro area or between countries. At that distance, copper was never a candidate — light has carried this traffic for decades, the same way it has carried telecom backbone traffic since the 1990s. This zone runs on coherent optics, a more sophisticated and more expensive technology than the pluggable modules used inside a data hall, built to survive far longer fiber runs without the signal needing to be regenerated as often.

Exhibit 5: World map of underwater Fiber optic cable (Scale-Across)



Sources: wbur.org

Scale-Out — where the market actually is today

Scale-Out connects racks to each other within a data hall — and the reason it made the jump from copper to light years ago, comes down to the same trade-off previous section introduced: **the faster a copper link runs, the shorter it must be.**

At 400 Gbps per lane, already on the roadmap, copper's reach shrinks to centimeters. No better copper, thicker gauge, or cleverer cable design removes the underlying physics — it only delays the wall by a generation, which is exactly why Scale-Out moved first. Today this zone runs mostly at 800G, and is in the early stages of shifting to 1.6T. That speed and reach is carried end to end by a single piece of hardware: **the optical transceiver.**

An optical transceiver is a small pluggable module that plugs into the front panel of a switch, the same physical idea as plugging a USB drive into the front of a computer. Inside sits a laser (the light source), a **modulator** (encodes data onto the light), a **photodetector** (converts light back into an electrical signal at the receiving end), a driver/TIA chip (amplifies the signal), and a **DSP** chip (cleans up the signal on both sides of the optical hop). The whole assembly is hot-swappable — pulled out and replaced without taking the switch down.

Exhibit 6: Copper's shrinking reach as per-lane speed rises

Per-lane speed	Typical port	Practical copper reach	Note
100 Gbps/lane	800G	~2 meters	Spans most of a rack
200–224 Gbps/lane	1.6T (today's leading edge)	<1m passive; 2–3m only with an AEC (Active Electrical Cable — adds power-hungry chips just to boost the signal)	Reach is now a real design constraint
400 Gbps/lane	Next generation	A few centimeters	Effectively unusable outside one chip package

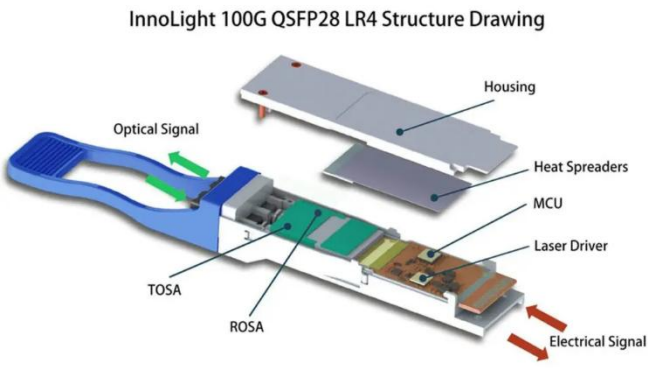
Sources: Globlex research

Exhibit 7: Pluggable optical transceiver



Sources: [my.avnet.com](https://www.avnet.com)

Exhibit 8: Pluggable optical transceiver components



Sources: [Keysight.com](https://www.keysight.com)

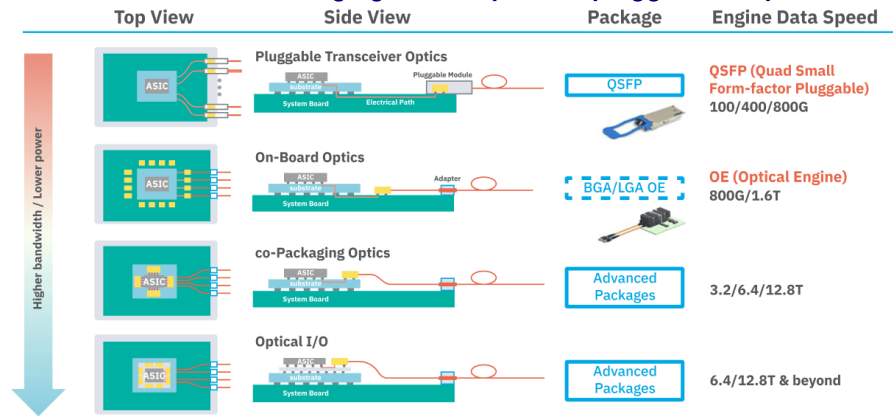
The transceiver has two structural problems the industry is working through.

- **The first is speed and power:** the DSP chip that cleans up the signal is expensive in both cost and power and works harder as more data is pushed through it.
- **The second is Indium Phosphide (InP) itself** — the laser inside a traditional transceiver is built entirely on InP, a difficult material to manufacture at low yield and rising cost.

The industry's near-term answer is **Silicon Photonics (SiPh)** replacing the traditional EML (Electro-absorption Modulated Laser, a design that combines the laser and its modulator in one InP chip — an 800G module might use eight of these). **SiPh is already in production and shipping today**, not a future technology: it keeps a single steady CW laser (continuous wave — an always-on beam) and hands the modulation job to a separate silicon chip instead, cutting laser-chip count from roughly eight down to one.

The next step beyond that is moving the optical engine physically closer to the switch or accelerator chip itself — **NPO** (Near-Packaged Optics, onto the same board as the switch chip) and **CPO** (Co-Packaged Optics, mounted on the very same package) — which is still in qualification rather than volume production today. Shorter electrical distance means less signal to clean up, so the DSP problem two paragraphs up shrinks with it: roughly 15 pJ/bit for pluggable down to ~5 pJ/bit for CPO

Exhibit 9: Photonic Packaging Roadmap: From pluggable to optical I/O



Sources: [Trendforce.com](https://www.trendforce.com)

Exhibit 10: Four Generations of optical integration — Energy cost per bit

Generation	What changes	Energy cost	Optical engine sits
Pluggable (today's standard)	Self-contained module plugs into the switch faceplate	~15 pJ/bit	Several cm from the chip
LPO / NPO	Removes some cleanup electronics, or moves onto the switch board	~8 pJ/bit	On-board, closer
CPO	Mounted on the same package as the switch chip	~5 pJ/bit	Same package
Optical I/O	Optical connectivity built into the compute die itself	~3 pJ/bit	Inside the chip package

Sources: Globlex research

Scale-Up

Scale-Up — GPU to GPU inside the same rack via NVLink: it's the zone set to face the fastest-rising bandwidth demand of the three, and the one zone still mostly running on copper today.

At Rubin, still 72 GPUs to a rack (NVL72), copper continues to hold — GPUs are close enough together that the wall described earlier in this paper hasn't been reached yet.

Rubin Ultra changes that. Its NVL576 domain scales the same idea across eight racks — 576 GPUs that still need to talk all-to-all, instantly, so the whole domain behaves as one giant GPU, the same design goal NVLink has always had, just stretched across a distance copper genuinely cannot bridge. Past that point, Scale-Up needs optical connectivity too — and, like Scale-Out before it, that means leaning on **CPO/NPO packaging** to keep the electrical hop short enough to afford.

The stakes are rising further with how AI inference itself is being architected. **Mixture-of-Experts (MoE)** designs treat each GPU less like an interchangeable unit of compute and more like a specialist — an "expert" handling one slice of the model — which means GPUs now need to exchange information with each other more often, not less, as inference scales. That makes Scale-Up's bandwidth problem more urgent, not something the industry can defer to the next generation.

Exhibit 11: NVLink Bandwidth by GPU Generation, 2016–2027 (Scale-Up)

GPU Generation	NVLink Version	Total Bandwidth per GPU (bidirectional)	Bandwidth per Link (bidirectional)	Number of Links	Year
Pascal (P100)	NVLink 1.0	160 GB/s	40 GB/s	4	2016
Volta (V100)	NVLink 2.0	300 GB/s	50 GB/s	6	2017
Ampere (A100)	NVLink 3.0	600 GB/s	50 GB/s	12	2020
Hopper (H100/H200)	NVLink 4.0	900 GB/s	50 GB/s	18	2022
Blackwell (B200/GB200/GB300)	NVLink 5.0	1,800 GB/s	100 GB/s	18	2024–25
Rubin	NVLink 6.0	3,600 GB/s	200 GB/s (not official)	18 (not official)	H2 2026
Rubin Ultra	NVLink 7.0	~2× Rubin (unconfirmed)*	TBD	TBD	H2 2027

Sources: NVIDIA; Globlex research

Sizing the Opportunity

Independent forecasts differ in scope, definition and time horizon, but they agree on direction: this is a genuinely fast-growing market, on both a one-year view and out to the end of the decade.

Exhibit 12: Optical connectivity market size forecasted by source

Source	Period	Initial Market Size	End-of-Period Size	CAGR / YoY Growth
LightCounting	2025–2026	USD 16.50bn	26.00bn	60.00%
Yole Group	2025–2031	USD 23.40bn	112.30bn	30.00%
CIC	2025–2030	USD 24.80bn	111.00bn	31.60%

Sources: Alphasense

The near-term figure and the longer-horizon ones are measuring somewhat different scopes of the market, so they shouldn't be read as three estimates of the same number — but the pattern holds across all three: LightCounting has the market up ~60% in the year immediately ahead alone (2025→2026), while Yole Group and CIC, on longer 5-6 year windows, both converge on the market compounding to roughly US\$111-112 billion by 2030-2031 — close to a 4.5-5x expansion from today's base, at CAGRs in the 30-32% range.

It's worth thinking about this growth as rising optical content per unit of compute — the dollars of optical hardware a data center buys per GPU, or per unit of AI compute deployed — since that ratio, not just GPU unit growth, is what's actually climbing as networks get faster and denser. Content climbs from ~120-130 modules at H100 to ~432 by Vera Rubin — all Scale-Out, since Scale-Up still runs on copper within one rack. Rubin Ultra changes the topology: eight racks merge into one Scale-Up domain, a distance only CPO/NPO can bridge, so its ~80-100 optical engines aren't directly comparable to earlier counts. Cost is the reliable measure instead — from ~USD25-35K at H100 to an estimated USD1.1-1.2 million by Rubin Ultra, a ~44x increase — making optics a steadily larger slice of the rack's total bill of materials with each generation that passes.

Exhibit 13: Optical content, components and costs, per rack by Generation

GPU Generation	Architecture	Est. Optical Content / Rack	Est. Optical Cost / Rack
Hopper	H100 / H200	~120–130 modules	USD25–35K
Blackwell	GB200 NVL72	~180–250 modules	USD100–150K
Blackwell Ultra	GB300 NVL72	~216 modules	USD315K
Vera Rubin	Rubin NVL72	~432 modules	USD490–500K
Rubin Ultra	Rubin NVL576	~80–100 optical engines*	USD1.1–1.2M

Sources: Alphasense; Globlex research

Three forces behind the growth

- **First, unit growth.** More AI clusters simply means more transceivers, fiber, and switches sold, independent of any technology shift.
- **Second, speed migration.** Each new port generation (800G → 1.6T → beyond) commands a materially higher price than the one it replaces, because it packs more optical and electrical complexity into the same footprint.
- **Third, technology-mix shift.** As CPO moves from pilot to volume, content per port rises further, because it requires the harder, semiconductor-grade packaging—not the simpler optical assembly pluggable modules use today.

Together, these three forces compound rather than simply add up — unit growth expands how many racks get built, while speed migration and technology-mix shift raise the dollar content within each one. That's why optical spending keeps outpacing GPU unit growth alone. The remaining question is who actually captures that growing dollar

Who benefits: A Layer-by-Layer map

Photonics doesn't compress into two or three names the way the GPU market clusters around NVIDIA, or HBM around a handful of memory makers. It's a genuinely fragmented chain — whether the finished product is today's pluggable transceiver or tomorrow's NPO/CPO architecture, who actually captures the value is decided by many companies working at very different stages, not one or two. To make that legible, we split the photonics supply chain into seven genuinely different layers, each with its own materials, factories and small set of key beneficiary companies.

Layer 1 — Substrate: the raw material everyone needs. Every laser starts as a raw InP (Indium Phosphide) wafer, grown from indium — a metal that exists only as a byproduct of zinc smelting, not something a dedicated mine can simply produce more of.

Key beneficiaries: AXT, Sumitomo Electric, JX Advanced Metals.

Layer 2 — Light source: turning the wafer into a laser. Silicon has an indirect bandgap and can't efficiently emit light; InP has a direct bandgap and can. Turning a bare InP wafer into a working laser chip — epitaxial growth, then device fabrication — is a separate step that the companies who do it keep in-house rather than buy in.

Key beneficiaries: Coherent, Lumentum, Mitsubishi Electric.

Layer 3 — Electrical IC: the signal side. Light doesn't turn itself on. An electrical signal has to drive the laser or modulator at the sending end, and a receive-side amplifier (a TIA) must turn a faint returning light signal back into a readable voltage. That conversion never disappears, whatever the module design — pluggable, LPO, or CPO.

Key beneficiaries: Broadcom, Marvell, Credo Technology.

Layer 4 — Photonic chip foundry: the manufacturing platform. A photonic integrated circuit (PIC) — the chip that guides, splits and modulates light but can't generate it — is built on largely the same equipment as ordinary chips, at mature process nodes. Increasingly, foundries here also offer the advanced packaging platform CPO needs to physically bind an optical chip to an electrical one.

Key beneficiaries: TSMC, GlobalFoundries, Tower Semiconductor.

Layer 5 — Module & packaging assembly: where it all comes together — and where CPO changes who benefits. This is where the laser die, photonic die and electrical die are physically joined into one working transceiver. For today's pluggable modules, that's a contained optical-assembly step, and the winning skill is sourcing chips cheaply, assembling at high yield, and requalifying fast each time

speeds step up — the game Innolight and Eoptolink have mastered. CPO raises the bar to a semiconductor-grade **OSAT** packaging job, and, per Layer 4 above, increasingly hands the hardest part of that job to the foundry itself, leaving this layer a harder board- and system-level integration role.

Key beneficiaries: Innolight, Eoptolink, Fabrinet.

Layer 6 — Coherent, DCI & fiber: Two things sit here: the coherent-optics gear that carries Scale-Across traffic between data centers, and the fiber-optic cable itself — drawn from silica glass, entirely outside the chip chain above it, yet just as essential.

Key beneficiaries: Ciena, Nokia, Corning.

Layer 7 — System: who decides the architecture. Hyperscalers and GPU vendors sit at the top, and they're not just customers — they decide whether a given link stays copper or goes optical, and whether the ecosystem below them is proprietary or open.

Key beneficiaries: NVIDIA, Broadcom, Google.

Exhibit 143: The seven-layer value chain of photonics

#	Layer	What it does	Key names
1	Substrate	Grows the raw InP wafer every laser starts from	AXT, Sumitomo Electric, JX Advanced Metals
2	Light source	Turns the InP wafer into a working laser chip	Coherent, Lumentum, Mitsubishi Electric
3	Electrical IC	Drives the laser and reads the signal back — fabless, manufactured at Layer 4	Credo Technology, Broadcom, Marvell
4	Photonic foundry	Fabricates the PIC and Layer 3's IC; increasingly does CPO's die-level packaging too	TSMC, GlobalFoundries, Tower Semiconductor
5	Module & packaging assembly	Join laser, PIC and electrical dies into a finished module	Innolight, Eoptolink · Fabrinet
6	Coherent, DCI & fiber	Carries Scale-Across traffic; supplies the fiber Layer 5 needs too	Ciena, Nokia, Corning
7	System	Decides the architecture — copper or light, proprietary or open	NVIDIA, Broadcom, Google

Sources: Globlex research

Investment Implications

Rather than profiling each company in isolation, it's more useful to sort them by which technology clock they're actually riding — because that's what decides whether today's setup helps a name now, still to come, or both.

Near-term: winning while pluggable holds the line

Benefits the longer 800G/1.6T pluggable stays dominant — still most of the market today, and expected to coexist with CPO for a few more years.

Innolight and Eoptolink (Layer 5) — the pluggable-era assembly leaders: cheap sourcing, high yield, fast requalification at every speed grade. Every extra year pluggable stays dominant extends their market directly.

Credo Technology (Layer 3) — makes the DSP and AEC chips that let copper stretch a little further. Every meter photonics hasn't replaced yet is still Credo's market.

Medium-to-longer term: winning once NPO/CPO reaches volume

The reward arrives once CPO moves from qualification to real shipping volume — priced on a longer runway than the names above.

TSMC (Layer 4) — betting CPO packaging becomes a foundry business, not just a fabrication one. Its COUPE platform does the die-level packaging that used to sit entirely with Layer 5.

Broadcom and Marvell (Layer 3 & 7) — gain as the "transceiver business" migrates into the switch package itself: switch-ASIC and custom-silicon vendors capture value that used to sit with module assemblers.

Both legs: needed whichever way it goes

Required by every generation, not just one — so timing barely matters.

AXT and the InP substrate suppliers (Layer 1) — every laser, in every module form, still starts as an InP wafer. Worth flagging: CPO's higher-power CW laser tends to mean lower yield per wafer, so the shift to CPO may raise InP demand rather than shrink it. AXT itself supplies roughly a third of global InP, largely via its Chinese subsidiary Beijing Tongmei — a concentration that cuts both ways as export-licensing policy evolves.

Coherent and Lumentum (Layer 2) — the two vertically-integrated laser makers, taking substrate through epitaxy to a finished chip in-house, a capability that matters the same way in a pluggable module or a CPO package.

Corning (Layer 6) — fiber is needed no matter which architecture wins; the one genuinely form-neutral layer in this chain.

NVIDIA's own capital backs this reading: Coherent and Lumentum each received a US\$2 billion investment from NVIDIA, while Corning received roughly US\$500 million in warrants — funding aimed at capacity that serves pluggable and NPO/CPO production alike, reinforcing our view that Layer 2 and Layer 6 collect value in both the pluggable-transceiver era and the NPO/CPO era to come.

The shift from copper to light is not a speculative bet — it's a consequence of physics that gets more binding, not less, with every new GPU generation. What remains open is the pace: how quickly CPO moves from qualification to volume, how the EML-versus-CW-laser design choice settles, and how durable today's supply concentration proves under commercial and geopolitical pressure.

GENERAL DISCLAIMER

Analyst Certification

Suwat Sinsadok, Register No. 020799, Globlex Securities Public Company Limited

The opinions and information presented in this report are those of the Globlex Securities Co. Ltd. Research Department. No representation or warranty in any form regarding the accuracy, completeness, correctness or fairness of opinions and information of this report is offered by Globlex Securities Co. Ltd. Globlex Securities Co. Ltd. Accepts no liability whatsoever for any loss arising from the use of this report or its contents. This report (in whole or in part) may not be reproduced or published without the express permission of Globlex Securities Co. Ltd.

RECOMMENDATION STRUCTURE

Stock Recommendations

Stock ratings are based on absolute upside or downside, which we define as $(\text{target price}^* - \text{current price}) / \text{current price}$.

BUY: Expected return of 10% or more over the next 12 months.
HOLD: Expected return between -10% and 10% over the next 12 months.
REDUCE: Expected return of -10% or worse over the next 12 months.

Unless otherwise specified, these recommendations are set with a 12-month horizon. Thus, it is possible that future price volatility may cause temporary mismatch between upside/downside for a stock based on market price and the formal recommendation.

* In most cases, the target price will equal the analyst's assessment of the current fair value of the stock. However, if the analyst doesn't think the market will reassess the stock over the specified time horizon due to a lack of events or catalysts, then the target price may differ from fair value. In most cases, therefore, our recommendation is an assessment of the mismatch between current market price and our assessment of current fair value.

Sector Recommendations

Overweight: The industry is expected to outperform the relevant primary market index over the next 12 months.
Neutral: The industry is expected to perform in line with the relevant primary market index over the next 12 months.
Underweight: The industry is expected to underperform the relevant primary market index over the next 12 months.

Country (Strategy) Recommendations

Overweight: Over the next 12 months, the analyst expects the market to score positively on two or more of the criteria used to determine market recommendations: index returns relative to the regional benchmark, index sharpe ratio relative to the regional benchmark and index returns relative to the market cost of equity.

Neutral: Over the next 12 months, the analyst expects the market to score positively on one of the criteria used to determine market recommendations: index returns relative to the regional benchmark, index sharpe ratio relative to the regional benchmark and index returns relative to the market cost of equity.

Underweight: Over the next 12 months, the analyst does not expect the market to score positively on any of the criteria used to determine market recommendations: index returns relative to the regional benchmark, index sharpe ratio relative to the regional benchmark and index returns relative to the market cost of equity.