

AI Infrastructure Primer

The AI Infrastructure Stack: A Deep Dive Primer for Investors

Understanding How Your Holdings Power the AI Revolution

Executive Summary

The AI revolution is reshaping data center infrastructure from the ground up. Traditional server-centric architectures are being replaced by rack-scale compute systems that require fundamentally different approaches to power, cooling, connectivity, and networking. This primer explains how your portfolio companies—Aster Labs (ALAB), Credo Technology (CRDO), Fabrinet (FN), nVent Electric (NVT), Tower Semiconductor (TSEM), Arista Networks (ANET), Coherent (COHR), and Quanta Services (PWR)—are positioned in this transformation.

The key insight: AI infrastructure isn't just about bigger servers. It's about reimagining the entire stack from silicon photonics to liquid cooling, where "boring" infrastructure companies often have the most defensible margins.

1. The AI Infrastructure Stack (Bottom to Top)

Power & Cooling: The Foundation Layer (PWR, NVT Territory)

Why AI Changes Everything

Traditional enterprise servers consume 150-300W per socket. Modern AI accelerators like NVIDIA's H200 consume up to 700W per GPU, with training clusters requiring 8 GPUs per node. A single AI training server can consume 6-8kW, compared to 1-2kW for traditional servers. Scale this to 100,000+ GPU clusters, and you're looking at 100-500MW data centers—equivalent to small cities.

The Physics Problem:

- **Compute Density:** AI chips pack more transistors per square mm, generating heat proportional to power consumption
- **Heat Removal:** Air cooling hits physical limits around 250W/cm². Beyond this, liquid cooling becomes mandatory
- **Grid Constraints:** Traditional data centers consume 10-30MW. AI data centers need 100-1000MW, straining electrical grids

nVent Electric (NVT): The Boring Infrastructure Play

nVent specializes in the unglamorous but critical infrastructure:

Liquid Cooling Systems:

- **Direct-to-chip cooling:** Liquid cooling blocks mounted directly on GPUs
- **Immersion cooling:** Entire servers submerged in dielectric fluid
- **Cold plates and heat exchangers:** Custom thermal management solutions

Power Distribution:

- **Intelligent PDUs:** Power distribution units with remote monitoring
- **Busway systems:** High-capacity power distribution for megawatt deployments
- **Thermal management:** Cable management systems optimized for liquid cooling

Recent Developments:

- Partnership with Siemens for NVIDIA AI data center reference architectures
- Modular liquid cooling solutions aligned with chip manufacturers' thermal requirements

- Focus on 400V/800V DC power distribution for efficiency

Quanta Services (PWR): The Execution Layer

Quanta builds the physical infrastructure:

- **Electrical grid connections:** Substations, transmission lines for massive power draws
- **Data center construction:** Specialized electrical and mechanical systems
- **Emergency power systems:** Backup generators, UPS systems for critical AI workloads

Investment Thesis: AI data centers require 10-50x more electrical infrastructure than traditional facilities. Quanta benefits from both new construction and grid upgrades.

Semiconductor Manufacturing: The Silicon Foundation (TSEM Territory)

Tower Semiconductor (TSEM): Beyond Leading Edge

While most investor attention focuses on leading-edge nodes (3nm, 5nm), AI infrastructure requires diverse silicon technologies:

Specialty Foundry Focus:

- **Analog/Mixed-Signal:** Power management ICs, SerDes drivers, clock distribution
- **Silicon Photonics:** Optical transceivers, modulators, photodetectors
- **SiGe (Silicon-Germanium):** High-frequency RF components, optical receivers
- **Power Electronics:** GaN (Gallium Nitride) and SiC (Silicon Carbide) for efficient power conversion

Why Specialty Matters:

Leading-edge logic (GPUs, CPUs) gets headlines, but AI infrastructure requires thousands of support chips:

- Power management ICs for clean power delivery
- SerDes transceivers for high-speed I/O
- Optical components for data center networking
- RF chips for wireless backhaul

Tower's Position:

- **65nm-180nm processes:** Optimized for analog performance, not transistor density
 - **Silicon photonics capabilities:** Critical for optical interconnects
 - **Automotive qualification:** Reliability standards applicable to data center infrastructure
-

Optical Interconnects: Breaking the Bandwidth Wall (COHR, FN Territory)

The Fundamental Problem

Electrical interconnects hit physical limits around 56Gbps per lane due to:

- **Skin effect:** High-frequency signals travel on copper surface, increasing resistance
- **Crosstalk:** Adjacent wires interfere at high frequencies
- **Power consumption:** Electrical SerDes consume exponentially more power at higher speeds

Optical interconnects use photons instead of electrons, enabling:

- **Higher bandwidth:** 100Gbps+ per optical lane
- **Lower power:** Photons don't have resistance
- **Longer reach:** Kilometers vs. meters for electrical

Coherent (COHR): The Vertical Integration Play

Coherent's strength lies in vertical integration across the optical stack:

Silicon Photonics:

- **Integrated photonic circuits:** Lasers, modulators, detectors on single chip
- **1.6T transceivers:** 1.6 terabit/second modules using Marvell DSPs
- **Co-packaged optics:** Optical components integrated directly with ASICs

Manufacturing Capability:

- **VCSEL arrays:** Vertical Cavity Surface Emitting Lasers for short-reach links
- **Optical components:** Lenses, filters, couplers, isolators
- **Assembly and test:** Complete transceiver manufacturing

Product Portfolio:

- **DR8 modules:** 8 lanes × 200Gbps = 1.6Tbps total bandwidth
- **Optical circuit switches:** Direct optical connections between racks
- **Active optical cables:** Integrated transceivers and fiber

Fabrinet (FN): The Manufacturing Partner

Fabrinet provides precision manufacturing for optical components:

- **Optical transceiver assembly:** Sub-micron precision alignment
- **Electro-optical PCBA:** Mixed analog/digital boards for transceivers
- **Supply chain management:** Southeast Asia manufacturing cost advantages

Investment Angle: As optical complexity increases, specialized manufacturing becomes more valuable. Fabrinet benefits from the shift to silicon photonics and integrated optics.

High-Speed Connectivity: The SerDes Bottleneck (CRDO Territory)

What is SerDes?

SerDes (Serializer/Deserializer) chips convert parallel data streams to high-speed serial links. They're the unsung heroes of AI infrastructure, enabling:

- **PCIe connections:** CPU to GPU communication
- **Ethernet:** Server-to-switch networking
- **Intra-rack links:** GPU-to-GPU communication

Credo Technology (CRDO): The SerDes Specialist

Credo's core innovation is making electrical SerDes work at higher speeds through:

Advanced Signal Processing:

- **Equalization:** Correcting signal distortion from cables/traces
- **Error correction:** Forward error correction (FEC) for reliable transmission
- **Retiming:** Cleaning up signal timing for longer reaches

Product Portfolio:

- **224Gbps retimers:** Enabling longer PCIe Gen6 and Ethernet connections
- **AEC (Active Electrical Cables):** Integrated SerDes in cable assemblies
- **Multiprotocol support:** UALink, Ethernet, PCIe on single chip

The AI Differentiator:

- **HiWire AECs:** Active cables with integrated retimers for reliable high-speed links
- **ZeroFlap technology:** Reducing optical link failures in massive AI clusters
- **PILOT SDK:** Debug tools for system-level optimization

Investment Thesis: As AI clusters scale to 100,000+ GPUs, SerDes reliability becomes critical. Credo's system-level approach addresses the #1 problem in massive AI deployments: link failures.

Custom Silicon: Beyond Merchant Silicon (ALAB Territory)

The Custom vs. Merchant Decision

Merchant Silicon: Standard chips (Intel CPUs, AMD GPUs) used by multiple customers **Custom Silicon:** Chips designed for specific applications/customers

AI workloads drive custom silicon because:

- **Workload optimization:** Custom datapaths for transformer models
- **Integration:** Combining multiple functions on single chip
- **Differentiation:** Custom features not available in merchant silicon

Astera Labs (ALAB): The Connectivity ASIC Company

Unlike Broadcom/Marvell (general networking ASICs), Astera focuses specifically on AI connectivity:

Product Categories:

- **PCIe retimers:** Enabling GPU-CPU communication over longer distances
- **CXL (Compute Express Link) controllers:** Memory expansion for AI workloads
- **Ethernet ASICs:** Custom networking chips for AI fabrics

The "Rack-Scale" Vision:

Traditional servers are GPU-limited. Astera enables rack-scale computing where:

- **Disaggregated resources:** Separate compute, memory, storage across rack
- **Flexible interconnects:** Dynamic allocation of resources via software
- **Open standards:** CXL, PCIe, Ethernet instead of proprietary interconnects

Recent Developments:

- **NVLink Fusion:** Custom solutions for NVIDIA GPU interconnects
- **COSMOS software:** Management and optimization for large fleets
- **aiXscale acquisition:** Adding photonics capabilities

Investment Differentiation: While Broadcom focuses on switching ASICs, Astera targets the connection layer—retimers, controllers, and bridges that make rack-scale AI possible.

Networking Platforms: The AI Fabric (ANET Territory)

The Shift from InfiniBand to Ethernet

Traditional HPC used InfiniBand for low-latency GPU clustering. AI is driving a shift to Ethernet because:

- **Scale:** Ethernet switches scale to 100,000+ ports vs. InfiniBand's smaller fabrics
- **Open standards:** Multiple vendor ecosystem vs. single-vendor InfiniBand
- **Economics:** Ethernet benefits from hyperscale volume economics

Arista Networks (ANET): The AI Ethernet Leader

Arista's strength lies in software-defined networking optimized for AI:

AI Spine Architecture:

- **Spine-leaf topology:** Non-blocking communication between any two points
- **800G Ethernet:** Latest generation providing 800Gbps per port

- **AI-optimized routing:** Multipath load balancing for training traffic

EOS (Extensible Operating System):

- **Unified software:** Same OS across all switches
- **Programmability:** Custom routing protocols for AI workloads
- **Telemetry:** Deep visibility into AI training job performance

Competitive Advantages:

- **Software differentiation:** EOS enables rapid feature development
- **AI-native design:** Purpose-built for AI training traffic patterns
- **Ecosystem partnerships:** Integration with NVIDIA, Astera Labs, and others

Technology Roadmap:

- **1.6T Ethernet:** Next generation doubling bandwidth again
- **Co-packaged optics:** Integrated photonics for lower power/cost
- **AI traffic optimization:** Machine learning-driven congestion management

2. The Data Flow: Tracing an AI Training Job

Let's trace a single AI training batch through the infrastructure stack:

Step 1: Power Delivery (PWR/NVT)

Grid Power → Substation (PWR) → Data Center PDU (NVT) → Server PSU → GPU (700W)

- Quanta builds grid connections for 100+ MW delivery
- nVent PDUs distribute and monitor power with 99.999% uptime
- Liquid cooling removes heat (nVent cold plates)

Step 2: GPU Computation

CPU sends training batch → PCIe (CRD0 retimer) → GPU memory → Compute → Results

- Credo retimers enable reliable PCIe Gen5/6 over longer traces
- Custom power management (TSEM ICs) maintains clean power delivery

Step 3: GPU-to-GPU Communication (ALAB)

GPU A → NVLink → Astera retimer → PCIe switch → NVLink → GPU B

- Astera's NVLink Fusion chips enable flexible GPU interconnection
- Data flows at 900GB/s between GPUs in same server

Step 4: Rack-to-Rack Communication (ANET/COHR)

Server → Ethernet (ANET switch) → Optical (COHR transceiver) → Fabric → Remote GPU

- Arista spine-leaf fabric routes training data between racks
- Coherent 800G/1.6T transceivers carry data over fiber
- AI training synchronization requires sub-microsecond latency

Step 5: Storage Access

GPU → PCIe → NVMe SSD → Training data → Memory → GPU processing

- High-bandwidth storage for dataset loading

- CXL memory expansion (Astera controllers) for larger models

Step 6: Manufacturing Integration (FN/TSEM)

Silicon design (TSEM foundry) → Assembly (FN) → Integration → Deployment

- Tower manufactures analog/mixed-signal support chips
- Fabrinet assembles complex optical transceivers

3. Key Concepts Investors Need to Know

Wafer Starts, Yields, and Advanced Packaging

Traditional Metrics Don't Apply

Wafer Starts: Number of silicon wafers entering production

- **Leading-edge focus:** 3nm/5nm wafers get attention but represent <10% of industry volume
- **AI infrastructure reality:** Most chips use 28nm-180nm processes optimized for specific functions

Yield Curves:

- **Digital logic:** Yield improves predictably with maturity
- **Analog/RF:** Yield limited by process variation, not defect density
- **Silicon photonics:** Optical alignment adds manufacturing complexity

Advanced Packaging Revolution:

Traditional chips: Single die in package AI chips: Multiple dies connected with high-bandwidth interconnects

CoWoS (Chip-on-Wafer-on-Substrate):

- **2.5D packaging:** Multiple dies on silicon interposer
- **3D packaging:** Stacked memory on logic die
- **Bandwidth advantage:** 1000x bandwidth vs. traditional packaging

Investment Implication: Packaging becoming more important than node shrinking. Companies with advanced packaging capabilities (TSMC, ASE Group) capture increasing value.

The Bandwidth Wall: Why Compute Outgrows Interconnect

Moore's Law vs. Communication Physics

Compute scaling: Transistor density doubles every 2 years (Moore's Law) **Interconnect scaling:** Physical limits constrain bandwidth growth

The Math:

- GPU compute: 2x every 2 years
- Electrical I/O: 1.3x every 2 years
- Gap widens to 10x over decade

Optical Solution:

Photonics can scale bandwidth faster than electronics:

- **Wavelength division multiplexing:** Multiple colors on single fiber
- **Spatial multiplexing:** Multiple fibers per cable
- **Advanced modulation:** More bits per photon

Investment Angle: Companies solving the bandwidth wall (Coherent, Astera Labs) capture increasing value as the gap widens.

InfiniBand vs. Ethernet for AI

The Technology Trade-offs

InfiniBand Advantages:

- **Lower latency:** Hardware-optimized for HPC workloads
- **RDMA:** Direct memory access without CPU involvement
- **Proven ecosystem:** Decades of HPC optimization

Ethernet Advantages:

- **Scale:** Spine-leaf topologies support 100,000+ endpoints
- **Economics:** Hyperscale volume drives cost down
- **Open standards:** Multiple vendor ecosystem

The Shift Underway

Early AI (2020-2023): InfiniBand dominated AI training **Current trend (2024-2026):** Ethernet gaining share **Future (2027+):** Ethernet expected to dominate

Driving Forces:

- **Scale requirements:** 100,000+ GPU clusters favor Ethernet
- **Cloud architecture:** Hyperscalers prefer open standards
- **Cost pressure:** Ethernet economics beat InfiniBand at scale

Investment Impact:

- **Winners:** Ethernet specialists (Arista, Broadcom)
- **Challenged:** InfiniBand pure-plays (Mellanox, now part of NVIDIA)

Silicon Photonics vs. Traditional Transceivers

The Technology Evolution

Traditional Transceivers:

- **Separate components:** Laser, modulator, detector, electronics
- **Assembly-intensive:** Manual alignment, high manufacturing cost
- **Performance limits:** Speed/power trade-offs

Silicon Photonics:

- **Integrated approach:** All optical functions on single chip
- **CMOS-compatible:** Leverages semiconductor manufacturing scale
- **Performance scaling:** Higher speeds, lower power per bit

Manufacturing Advantage

Traditional approach:

Laser die + Modulator + Detector + Assembly → Transceiver
(Multiple suppliers, complex alignment)

Silicon photonics:

Photonic integrated circuit + Electronic IC + Simple assembly → Transceiver
(Single-chip integration, automated assembly)

Economic Model: Silicon photonics has higher R&D costs but lower manufacturing costs at volume. The crossover happens around 100,000 units/year—exactly where AI transceivers are heading.

Why “Boring” Infrastructure Has the Best Margins

The Defensibility Hierarchy

Most Defensible (Highest Margins):

1. **Physical infrastructure:** Power, cooling, enclosures (NVT, PWR)
2. **Custom silicon:** Application-specific chips (ALAB, CRDO)
3. **Specialized manufacturing:** Precision assembly (FN, TSEM)
4. **Software platforms:** Network OS, management (ANET)

Least Defensible (Commoditizing):

5. **Standard components:** Generic transceivers, switches
6. **Contract manufacturing:** High-volume, low-mix production

Why Infrastructure Wins

Switching costs: Redesigning power/cooling infrastructure takes years **Reliability requirements:** AI training jobs cost millions—uptime is everything **Regulatory moats:** Safety certifications, building codes create barriers **Scale economics:** Fixed costs amortized over large deployments

Example: nVent's Moat

- **Thermal engineering:** Decades of cooling optimization experience
- **Safety certifications:** UL, IEC approvals for data center deployment
- **Integration complexity:** Liquid cooling requires mechanical, electrical, and thermal expertise
- **Customer stickiness:** Once deployed, infrastructure changes are disruptive

CapEx Cycles: Who Benefits First/Last

The Hyperscaler Investment Cascade

Phase 1: Planning (12-18 months before deployment)

- **Winners:** Design/engineering services, power infrastructure (PWR)
- **Orders:** Grid connections, substations, data center construction

Phase 2: Infrastructure Build (6-12 months before)

- **Winners:** Physical infrastructure (NVT), semiconductor foundries (TSEM)
- **Orders:** Power distribution, cooling systems, silicon production

Phase 3: Equipment Assembly (3-6 months before)

- **Winners:** Manufacturing services (FN), component suppliers (COHR)
- **Orders:** Optical transceivers, cable assemblies, final integration

Phase 4: Deployment (0-3 months)

- **Winners:** Networking platforms (ANET), connectivity ASICs (ALAB, CRDO)
- **Orders:** Switches, retimers, software licenses

Phase 5: Operation (ongoing)

- **Winners:** Software platforms, support services
- **Revenue:** Subscription, maintenance, upgrades

Investment Strategy: Infrastructure companies (PWR, NVT) provide earliest signals of AI CapEx cycles. Equipment companies (ALAB, CRDO, ANET) show momentum 6-12 months later.

TAM Expansion: New Markets vs. Share Taking

Traditional Data Center vs. AI Infrastructure

Traditional Data Center TAM: ~\$50B annually

- **Servers:** Standard x86, modest power/cooling requirements
- **Networking:** 1G/10G Ethernet, traditional spine-leaf
- **Growth rate:** ~5-10% annually

AI Infrastructure TAM: ~\$200B+ by 2030

- **Specialized servers:** GPU-optimized, 5-10x power density
- **High-speed networking:** 400G/800G/1.6T, AI-optimized fabrics
- **Growth rate:** 30-50% annually

New Market Creation vs. Substitution

New Markets (TAM Expansion):

- **Liquid cooling systems** (NVT): Didn't exist at scale pre-AI
- **AI-specific retimers** (CRDO): New category for GPU interconnects
- **Silicon photonics transceivers** (COHR): Replaces traditional electrical I/O
- **Rack-scale connectivity** (ALAB): New architecture paradigm

Market Share Taking:

- **AI-optimized switches** (ANET): Taking share from legacy networking vendors
- **Specialty foundry** (TSEM): Growing faster than leading-edge foundries
- **Precision manufacturing** (FN): Taking share from traditional EMS providers

Investment Implication: Companies creating new markets (NVT liquid cooling, ALAB rack-scale) have higher growth potential than those taking share in existing markets.

4. Glossary

AEC (Active Electrical Cable): Cable assembly with integrated SerDes chips to extend electrical signaling reach and improve reliability.

ASIC (Application-Specific Integrated Circuit): Custom-designed semiconductor for specific application, offering better performance/power/cost than general-purpose chips.

Bandwidth Wall: The growing gap between compute performance growth and interconnect bandwidth growth, driving need for optical solutions.

CoWoS (Chip-on-Wafer-on-Substrate): Advanced packaging technology stacking multiple dies on silicon interposer for higher bandwidth.

CXL (Compute Express Link): Open standard protocol for high-speed CPU-to-device and CPU-to-memory connections, enabling memory/accelerator sharing.

DSP (Digital Signal Processor): Specialized processor optimized for digital signal processing, used in optical transceivers for signal recovery.

EOS (Extensible Operating System): Arista's network operating system providing unified management across switch platforms.

FEC (Forward Error Correction): Error correction technique adding redundancy to data transmission for reliable high-speed links.

InfiniBand: High-performance networking standard optimized for HPC/AI with hardware-accelerated RDMA capabilities.

Moore's Law: Observation that transistor density doubles approximately every two years, driving compute performance growth.

NVLink: NVIDIA's proprietary high-bandwidth interconnect for GPU-to-GPU communication in AI systems.

Optical Circuit Switch (OCS): All-optical switch providing direct fiber connections between endpoints without electrical conversion.

PDU (Power Distribution Unit): Equipment distributing electrical power within data centers, often with monitoring and control capabilities.

RDMA (Remote Direct Memory Access): Technology allowing direct memory access between servers without CPU involvement, reducing latency.

SerDes (Serializer/Deserializer): Circuits converting parallel data to high-speed serial streams for chip-to-chip communication.

SiGe (Silicon-Germanium): Semiconductor alloy combining silicon and germanium for high-frequency analog and RF applications.

Silicon Photonics: Technology integrating optical components with electronic circuits on silicon chips for high-speed data transmission.

Spine-Leaf Architecture: Non-blocking network topology where every leaf switch connects to every spine switch, enabling any-to-any communication.

TSMC CoWoS: Taiwan Semiconductor's advanced packaging technology enabling 2.5D and 3D chip integration.

UALink: Ultra Accelerator Link, open standard for high-speed interconnects between AI accelerators and systems.

VCSEL (Vertical Cavity Surface Emitting Laser): Type of semiconductor laser optimized for short-reach optical communications in data centers.

Wafer Start: Metric measuring number of silicon wafers entering semiconductor fabrication process.

Investment Conclusions

The AI infrastructure transformation creates a complex value chain where traditional server-centric thinking doesn't apply. Success requires understanding the physics constraints (bandwidth walls, thermal limits, power delivery) driving fundamental architecture changes.

Key Investment Themes:

1. **Infrastructure First:** Physical infrastructure companies (PWR, NVT) benefit earliest from AI CapEx cycles and have most defensible positions.
2. **Connectivity Bottlenecks:** SerDes and optical interconnect companies (CRDO, COHR, ALAB) solve the fundamental bandwidth scaling problem.
3. **Specialty Over Leading-Edge:** AI infrastructure needs diverse silicon technologies, benefiting specialty foundries (TSEM) over leading-edge pure-plays.
4. **Open Standards Win:** Ethernet-based solutions (ANET) gain share over proprietary alternatives as AI clusters scale.
5. **Manufacturing Complexity:** Precision manufacturing (FN) becomes more valuable as optical integration complexity increases.

Portfolio Positioning: The companies profiled represent different risk/reward profiles across the AI infrastructure stack. Infrastructure (PWR, NVT) offers defensive positioning with predictable growth. Connectivity specialists (ALAB, CRDO) provide higher growth potential with execution risk. Platform companies (ANET) balance growth and defensive moats through software differentiation.

The AI revolution is still in early innings, but the infrastructure requirements are well-defined by physics and economics. Understanding these constraints helps identify which companies will capture sustainable value as AI scales from thousands to millions of GPUs over the next decade.

This primer provides educational information for investment research purposes. Past performance does not guarantee future results. All investments carry risk of loss. Consult qualified financial advisors for investment decisions.