



Marvell Announces Availability of Industry's First 102.4 Tbps Switch Purpose-Built for AI and Cloud Data Center Infrastructure

- **Purpose-built for AI:** The Marvell® Teralynx® T100 delivers up to 25% lower power than competitive solutions and the industry's lowest latency for AI training and inference workloads, eliminating the legacy elements that add power and complexity.
- **Breaking the AI power wall:** As GPU racks approach 120KW, the advanced low-power T100 architecture enables data centers to deploy larger numbers of accelerators within existing power envelopes.
- **Highest radix and flatter fabrics:** Supports 512-port scale-out radix and advanced scale-up fabric protocols, enabling fewer network tiers and simpler architectures that reduce latency and TCO across large AI clusters.

SANTA CLARA, Calif.--(BUSINESS WIRE)-- [Marvell Technology, Inc.](https://www.marvell.com) (NASDAQ: MRVL), a leader in data infrastructure semiconductor solutions, today introduced Marvell Teralynx T100, the industry's first 102.4 Tbps switch silicon purpose-built for the AI era. Unlike legacy switching platforms designed for traditional enterprise and cloud data centers, the Teralynx T100 was architected from the ground up for AI—enabling the industry's lowest power consumption and lowest latency at this bandwidth tier to address critical bottlenecks in today's large AI clusters. The T100 will start sampling to customers beginning this quarter.

Breaking the AI Power Wall

Today's AI data centers face a critical power challenge, as network inefficiencies translate to underutilized GPUs and higher training costs. GPU- and XPU-based systems are approaching 120KW per rack, pushing air cooling to its limits and increasingly requiring complex and costly liquid cooling solutions. Switching and networking components consume approximately 15-25% of total rack power¹—making low-power switch silicon a strategic requirement.

At under 1000W typical power, the T100 delivers up to 25% lower power than competitive solutions, enabling data center operators to accelerate AI infrastructure build-outs by deploying significantly more accelerators within existing power envelopes without requiring additional power infrastructure.

Purpose-Built for AI

The Teralynx T100 eliminates the unnecessary legacy elements that increase power and die area in competing solutions through a monolithic 102.4 Tbps device on advanced 3nm process technology. This approach enables flatter, higher-radix fabrics optimized for demanding AI workloads by reducing the number of AI network tiers and optical links. High-

radix, high-bandwidth low-latency switches are key to increasing GPU utilization, lowering tail latencies and improving convergence times for training algorithms, and the T100 delivers exceptional bandwidth efficiency while lowering overall rack power consumption and improving cluster efficiency.

“As AI workloads evolve and scale exponentially, hyperscalers require network architectures that optimize latency, power and scalability simultaneously,” said Rishi Chugh, vice president and general manager, Data Center Switch Business Unit, at Marvell. “The Teralynx T100 was purpose-built for AI—designed without the legacy baggage that inflates power, and engineered to deliver the deterministic performance and efficiency required to scale next-generation data center infrastructure.”

Designed for Scale-Out and Scale-Up AI Fabrics

The Teralynx T100 is optimized for both scale-out and scale-up deployments. For scale-out deployments, the T100 supports up to a 512-port radix, enabling operators to consolidate network tiers, simplify architectures and reduce latency across large AI training clusters with tens of thousands of accelerators. For scale-up deployments, the product’s highly flexible and programmable pipeline architecture supports a variety of interconnect standards and emerging scale-up fabric protocols—including the Ethernet Scale-Up Networking (ESUN) protocol—as well as the latest Ultra Ethernet Consortium (UEC) requirements and evolving AI Ethernet fabrics.

“As hyperscalers expand AI clusters to tens of thousands of accelerators, data center infrastructure becomes a defining factor in network efficiency and performance,” said Alan Weckel, co-founder and technology analyst at 650 Group. “The Teralynx T100 architecture delivers significant advantages in latency, power efficiency, radix scalability and overall TCO—advantages that stem directly from its purpose-built AI design approach and are required for the continued scaling of data center infrastructure.”

Flexible Deployment Options and Advanced Traffic Management

With an advanced low-power SerDes architecture, the Teralynx T100 is available in multiple package configurations—including ball grid array (BGA), co-packaged copper (CPC) and co-packaged optics (CPO) implementations—giving hyperscalers and cloud operators the flexibility to deploy the T100 in various form factors with a broad variety of connectivity options. The T100 also delivers latency-optimized topologies, integrated telemetry, AI-native congestion control and proprietary traffic management logic required by advanced data center architectures.

High-Performance, Low-Latency Portfolio Legacy

The T100 builds on the industry-leading high-performance and low-latency legacy of the Teralynx product family. From 12.8 Tbps to 102.4 Tbps, the Teralynx portfolio delivers a full range of switching solutions for every tier of the data center. It also includes a software toolkit with an SDK, support for the Open Compute Project’s (OCP’s) Switch Abstraction Interface (SAI), and the Software for Open Networking in the Cloud (SONiC) operating system—enabling a broad, open ecosystem for faster infrastructure deployments.

Availability

The Marvell Teralynx T100 switch will begin sampling to customers this quarter.

About Marvell

To deliver the data infrastructure technology that connects the world, we're building solutions on the most powerful foundation: our partnerships with our customers. Trusted by the world's leading technology companies for over 30 years, we move, store, process and secure the world's data with semiconductor solutions designed for our customers' current needs and future ambitions. Through a process of deep collaboration and transparency, we're ultimately changing the way tomorrow's enterprise, cloud and carrier architectures transform—for the better.

Marvell and the M logo are trademarks of Marvell or its affiliates. Please visit www.marvell.com for a complete list of Marvell trademarks. Other names and brands may be claimed as the property of others.

This press release contains forward-looking statements within the meaning of the federal securities laws that involve risks and uncertainties. Forward-looking statements include, without limitation, any statement that may predict, forecast, indicate or imply future events, results or achievements. Actual events, results or achievements may differ materially from those contemplated in this press release. Forward-looking statements are only predictions and are subject to risks, uncertainties and assumptions that are difficult to predict, including those described in the "Risk Factors" section of our Annual Reports on Form 10-K, Quarterly Reports on Form 10-Q and other documents filed by us from time to time with the SEC. Forward-looking statements speak only as of the date they are made. Readers are cautioned not to put undue reliance on forward-looking statements, and no person assumes any obligation to update or revise any such forward-looking statements, whether as a result of new information, future events or otherwise.

¹ Based on SemiAnalysis data: <https://newsletter.semianalysis.com/p/100000-h100-clusters-power-network>; <https://newsletter.semianalysis.com/p/datacenter-anatomy-part-1-electrical>; <https://newsletter.semianalysis.com/p/ai-training-load-fluctuations-at-gigawatt-scale-risk-of-power-grid-blackout>

View source version on businesswire.com:
<https://www.businesswire.com/news/home/20260601564526/en/>

Media Contact:
George Millington
pr@marvell.com

Source: Marvell Technology, Inc.