



CXL Consortium Members – Statement of Support for CXL 4.0 Specification

Alibaba

“As a Promoter member of the CXL Consortium, Alibaba has been a strong advocate of the CXL ecosystem. We are excited about the release of CXL 4.0, which supports PCIe Gen 7.0 and port bundling, addressing the increasing memory bandwidth demands of modern cloud workloads. Furthermore, CXL 4.0 enhances memory reliability and serviceability through advanced error reporting and memory sparing capabilities. We believe CXL 4.0 represents another significant milestone in the evolution toward composable, scalable, and reliable rack-scale architectures for next-generation data centers.”

Jian Chen, Senior Director of Engineering, Alibaba Cloud

Alphawave Semi

Alphawave Semi proudly supports the latest version of Compute Express Link, CXL 4.0. CXL continues to offer best fit solutions for memory acceleration, effective management of large, disaggregated data pools, and data integration across heterogeneous accelerators. Its ability to guarantee bandwidth for latency-sensitive tasks while providing high throughput for bulk data gives it distinct advantages for next-generation AI and HPC applications. Alphawave Semi remains committed to advancing this standard, which will boost system performance and open up new possibilities for emerging AI, HPC, and related technologies.

David Kulansky, Product Marketing Director, Alphawave Semi

Arm

“The industry’s continued focus on standardization is critical to fostering innovation and interoperability across an increasingly diverse AI computing landscape. The CXL 4.0 specification represents another meaningful step toward enabling the scalable memory and compute architectures that will define next-generation AI centers.”

Dong Wei, Standards Architect and Fellow, Arm

Astera Labs

"The CXL 4.0 specification's doubling of bandwidth to 128 GT/s and addition of bundled ports directly addresses the memory bottleneck constraining today's AI infrastructure—enabling memory capacity and bandwidth to scale in tandem with exponential AI model demands. Additionally, as an open standard, CXL 4.0 delivers memory pooling and sharing capabilities for rack-scale AI systems, driving competitive innovation while providing the interoperability essential for multi-vendor infrastructure deployments."

Chris Petersen, Fellow, Technology & Ecosystems, Astera Labs and CXL Consortium Board Member



Intel

“Intel’s leadership in open standards continues to drive scalable, efficient architectures across AI, data center, and edge applications. CXL 4.0 is an exciting development to support the escalating performance demands of modern heterogeneous data-centric compute, including doubling of bandwidth with lower latency, bundled links, and RAS enhancements.”

Dr. Debendra Das Sharma, Intel Senior Fellow and Chair of CXL Board of Directors

Liquid

“The CXL 4.0 specification represents an important leap forward for the industry, doubling bandwidth while reinforcing the foundation for fully composable and uniquely scalable memory-centric infrastructure. As enterprises rely more on AI, HPC, and data-intensive workloads, this milestone will accelerate Liquid’s ability to deliver dynamic memory pooling and sharing solutions that dramatically improve performance, utilization, and sustainability across customer environments. We’re proud to contribute to the continued evolution of the CXL standard and the outcomes it enables.”

Sumit Puri, Co-Founder and Chief Strategy Officer, Liquid

Micron

“Micron is proud to support the release of the CXL 4.0 specification and played a pivotal role in defining new enterprise-class RAS features, including advanced memory error handling and standardized device maintenance capabilities like post-package repair (PPR). These critical enhancements deliver foundational reliability and serviceability requirements needed to deploy and manage large-scale, disaggregated memory pools with confidence. With speeds up to 128 GT/s and support for link aggregation, CXL 4.0 sets a new standard for performance. As a key contributor, Micron remains committed to advancing the CXL ecosystem and enabling the future of memory-centric computing.”

Siva Makineni, Vice President of the DRAM Systems Group, Micron Technology

Montage Technology

“As a contributing member of the CXL Consortium, Montage Technology is excited to support the release of the CXL 4.0 Specification. This milestone brings enhanced capabilities for scalable, disaggregated computing architectures, addressing critical memory challenges in data-intensive workloads. Leveraging our expertise in CXL memory expansion controller innovation, we look forward to collaborating with the consortium to accelerate CXL adoption and enable new levels of performance and efficiency for next-generation data centers.”

Stephen Tai, President, Montage Technology



Rambus

“Thanks to the leadership of the CXL Consortium, the industry continues to advance toward a computing architecture leveraging scalable, coherent, and pooled resources to meet the needs of AI and data-intensive workloads beyond the constraints of capacity or locality. The latest CXL IP advancements strengthen the fabric and resource pooling model, giving architects new flexibility to deliver performance, efficiency and resilience at rack scale.”

Matt Jones, SVP and GM of Silicon IP, Rambus

SK hynix

“SK hynix is pleased to support the release of the CXL4.0 specification that will accompany PCIe gen7 and Bundled Ports for higher bandwidth. We find that enabling CXL memory modules that can support PCIe gen7 speed will be critical for utilizing CXL solutions in AI use cases. By using CXL4.0 specification, SK hynix is looking forward to collaborate with the industry to enhance pooling and fabric related features that could allow to expand emerging applications and realize scalable memory disaggregation for the future.”

Uksong Kang, Vice President & Head of Next Gen. Product Planning and Enabling, SK Hynix

Synopsys

“As an active member of the CXL Consortium, Synopsys is committed to advancing the CXL 4.0 standard, which will enable scalable performance and unified memory fabric required for next-generation computing. Synopsys is already supporting the 128Gbps bandwidth that CXL 4.0 requires, with customers actively designing interoperable, high-bandwidth systems to scale AI infrastructure leveraging trusted IP solutions.”

Neeraj Paliwal, senior vice president of IP product management at Synopsys

UnifabriX

“UnifabriX welcomes the release of the new CXL 4.0 specification as a major milestone in advancing open, AI-scale fabrics, that interconnect CPUs, GPUs, accelerators, and memory, at rack level and beyond. CXL 4.0 is a major leap towards realizing next-generation AI infrastructures, doubling the per-lane communication rate to 128 GT/s and providing port aggregation for bandwidth-demanding applications, while enriching fabric capabilities, security, and resiliency. UnifabriX is excited to promote CXL into new applications and market segments by enabling transactional switching and Memory-over-Fabrics at near HBM-latencies, to power the most demanding inference and analytics workloads. The CXL Consortium continues to drive innovation at the forefront of interconnect technologies, and UnifabriX is proud to contribute to this transformative ecosystem.”

Ronen Hyatt, CEO and Chief Architect, UnifabriX