



CoreModule XL with Liquid SmartStack

Run AI workloads where and when required with composable GPU acceleration for Dell MX7000

Benefits and Features

High Performance GPU

- 30x 350W GPUs connected to the MX7000 chassis
- 8 Dual Xeon Compute sleds
- Native PCIe connection to GPUs in external chassis
- Large GPU memory capacity to run modern AI models

Compose resources on demand

- Based on Liquid Composable technology
- Configure CPU and GPU resources on-demand for each task
- Makes it easy to test and analyse and workload scaling
- Uses standard off-the-shelf GPUs and components

Learn More

EMEA Sales

+44 (0)20 7960 2400

emeasales@amulethotkey.com

N America Sales

+1 (212) 269 9300

ussales@amulethotkey.com

APJ Sales

+61 409 930 884

apsales@amulethotkey.com

AI Inference on-prem or at the Edge

CoreModule XL combined with Liquid SmartStack brings the power of composable GPU resource to the Modular MX7000 platform, enabling many more types of workload and protecting investment for the future.

Ideally suited for a wide range of AI training and inference workloads, Machine learning and Engineering Simulation, MX7000 with CoreModule XL makes it easy to leverage business-critical information in AI workflows whilst keeping everything secure and under control.

Deploy GPU resource where and when required

The composable architecture allows any number of GPU resources from the pool to be assigned to any compute sled on-demand, allowing workloads to be run in an optimal environment.



Flexibility for evolving needs

Designed for a range of compute requirements, the MX7000 platform can easily be configured for a range of applications and scales with your needs one GPU at a time.

Easily upgrade existing MX7000 deployments

CoreModule XL can be fitted to any existing MX7000 chassis using the Fabric A and B rear expansion bays to provide connectivity to the external GPU enclosure

www.amulethotkey.com/support

MX7000 CoreModule XL Brief - May 2024

