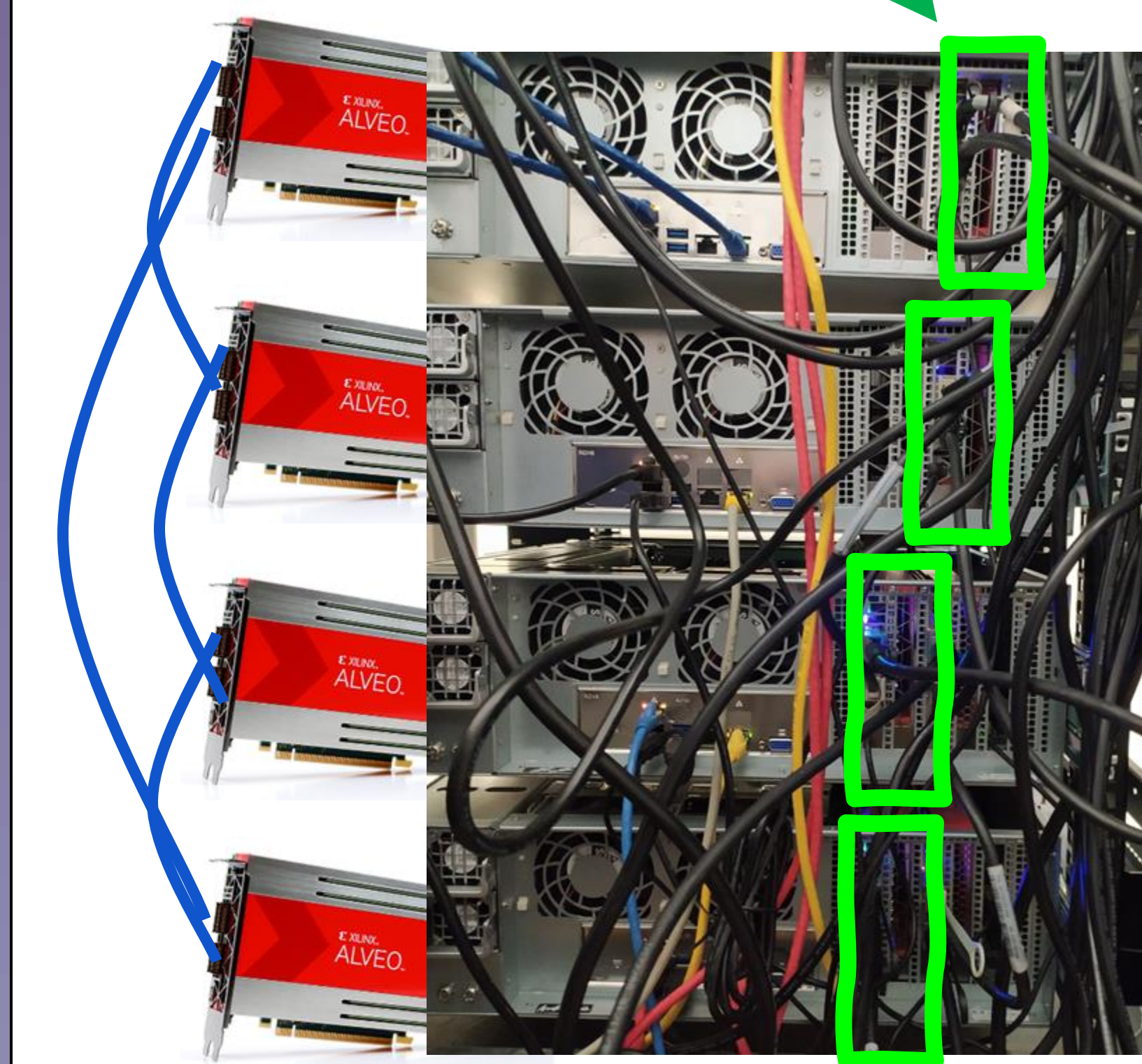


Figure 10 is a line graph titled "APEnetX (PCIe Gen3 X16) latency". The Y-axis is labeled "Time (us)" and ranges from 1.5 to 5.0. The X-axis is labeled "Message Size (Byte)" and is on a logarithmic scale with values 1, 2, 4, 8, 16, 32, 64, 128, 256, 512, and 1k. There are three data series:

- 2 hosts - synthetic ping-pong test (green line with circles):** This series shows the lowest latency for small message sizes, starting at approximately 2.0 us for 1 byte and decreasing slightly to about 1.9 us at 16 bytes, then increasing to about 2.1 us at 32 bytes and 2.3 us at 64 bytes. It then rises more steeply to about 2.7 us at 128 bytes, 3.7 us at 256 bytes, 5.0 us at 512 bytes, and 5.1 us at 1k bytes.
- 2 hosts - OSU latency benchmark (red line with squares):** This series shows slightly higher latency than the synthetic ping-pong test for small message sizes, starting at approximately 2.0 us for 1 byte and decreasing to about 1.9 us at 16 bytes, then increasing to about 2.1 us at 32 bytes and 2.3 us at 64 bytes. It then rises more steeply to about 2.7 us at 128 bytes, 3.7 us at 256 bytes, 5.0 us at 512 bytes, and 5.1 us at 1k bytes.
- Local loop - synthetic ping-pong test (blue line with triangles):** This series shows the lowest latency across all message sizes, starting at approximately 1.5 us for 1 byte and remaining constant until 32 bytes, then increasing to about 1.6 us at 64 bytes, 1.8 us at 128 bytes, 2.3 us at 256 bytes, 3.0 us at 512 bytes, and 4.7 us at 1k bytes.



- **Customizable I/O**
- **Deterministic latency**
- Useful in **heavy computation**

- INFN will contribute with **AI-Direct Engine** enabling the deployment of large scale NN models over multiple AIPU (Artificial Intelligence Processor Unit) accelerators to boost performance of applications like AI-accelerated HPC and Generative AI.
- **AI-Direct Specialized hardware** and its companion system software (linux device driver, user library) will be prototyped on a FPGA-based testbed

The diagram illustrates the system architecture, organized into several functional blocks and layers:

- HPC Runtimes & Applications:** Includes MPI, PGAS, and OFI. This section is highlighted with a red circle.
- Accelerators:** Includes AI Applications (e.g., PyTorch), NCCl/R, and RPC. This section is also highlighted with a red circle.
- Storage I/O:** Includes MPI-IO, Lustre, HDF5, Object (DAOS)/KV (Tebis), and NVMe-oF. This section is highlighted with a red circle.
- Datacenter Deployment:** Includes Datacenter services (e.g., distributed storage), Mercury RPC, Containers/VMs, and Sockets.
- Interfaces and Components:**
 - BXiv3 Portals:** Connects the HPC Runtimes & Applications, Accelerators, and Storage I/O sections.
 - User Space:** Includes Dashboard, Monitoring & Diagnosis, Configuration, Management, and Installation.
 - Kernel Space:** Includes IP/UDP/TCP and Ethernet.
 - NIC-level management, config, statistics:** Connects the User Space and Kernel Space.
 - GPU Direct:** Connects the HPC Runtimes & Applications and Accelerators.
 - Multi-NIC, NIC-triggered ops:** Connects the Storage I/O and Datacenter Deployment.
 - BXiv3 NIC:** Connects the Storage I/O and Datacenter Deployment.
 - Network Link/Switch:** Connects the Network Link/Switch and the Network Link/Switch.