

TPU Rack Overview



Pankaj Makhija, TPU Systems Architect & Tech Lead,
Google

August 24, 2025
Hot Chips Tutorial

Content

- 01 [TPU System Overview](#)
- 02 [Ironwood Rack](#)
- 03 [Challenges & Opportunities](#)

TPU System Overview

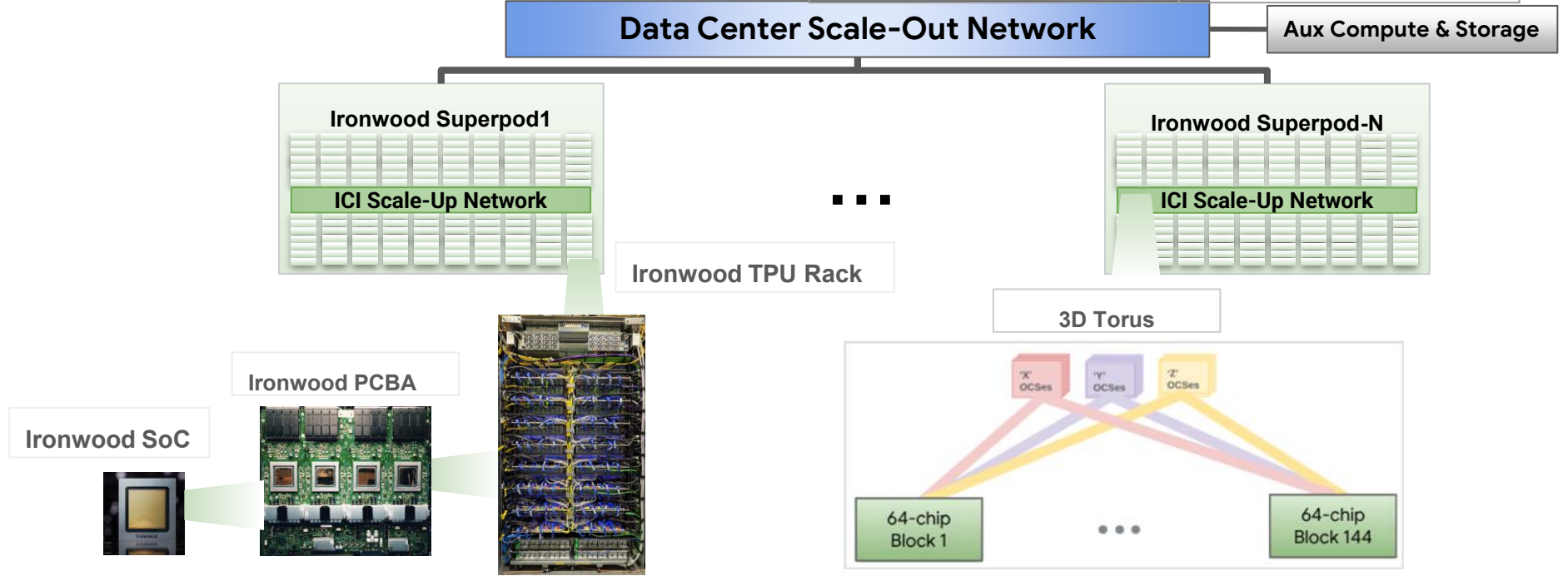
TPU Systems at a Glance

			
	TPU v4	TPU v5p	Ironwood
	2022	2023	2025
Pod Size (chips)	4096	8960	9216
HBM Bandwidth/ Capacity	32 GB @ 1.2 TBps HBM	95 GB @ 2.8 TBps HBM	192 GB @ 7.4 TBps HBM
Peak Flops per chip	275 TFLOPS	459 TFLOPS	4614 TFLOPS

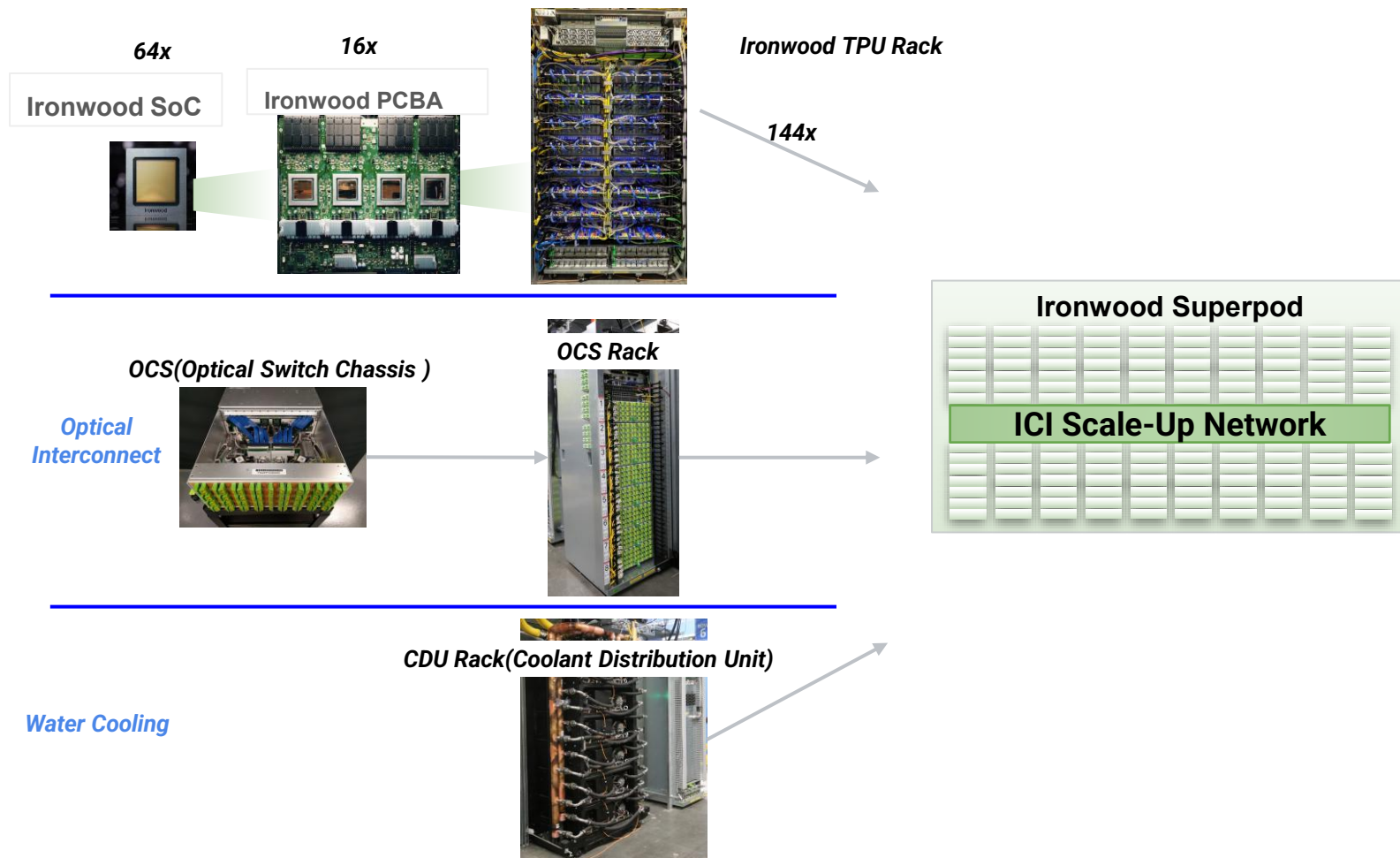
Ironwood Superpod & Max-scale Cluster

1.5x-10x Scaling across all the critical vectors: FLOPs, HBM BW, ICI BW

Traffic Type	Interface
Intra-Superpod Computation	ICI(Inter-Chip Interconnect)
Inter-Superpod Computation	NIC/s
Data Ingestion and Storage <=> Superpod	NIC/s

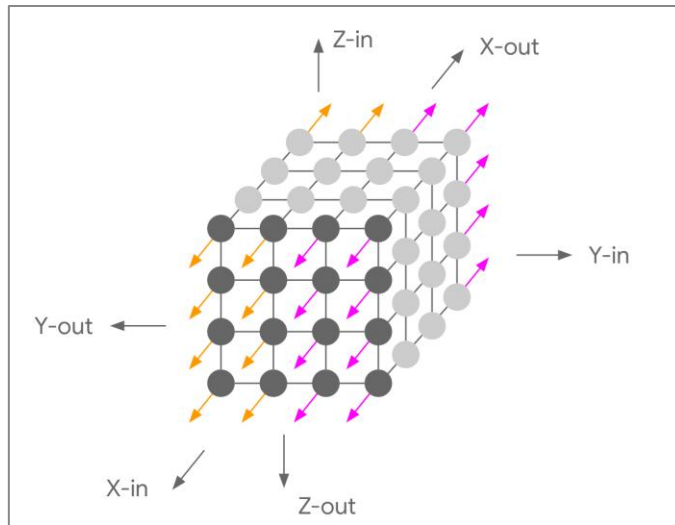
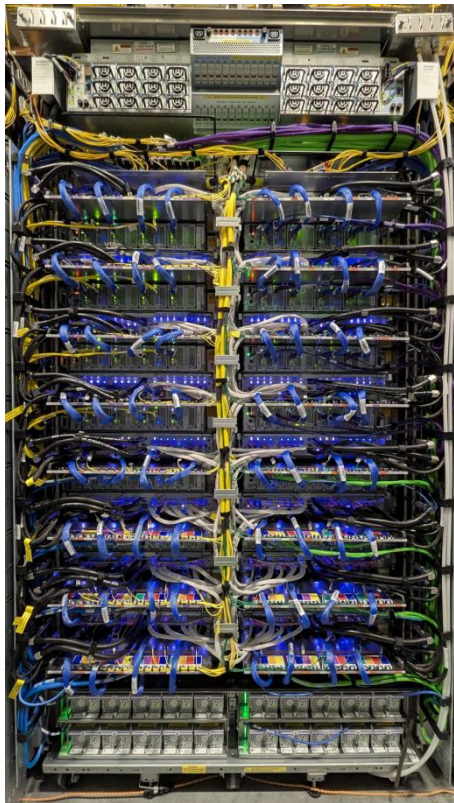


The Suite of Racks in a Superpod



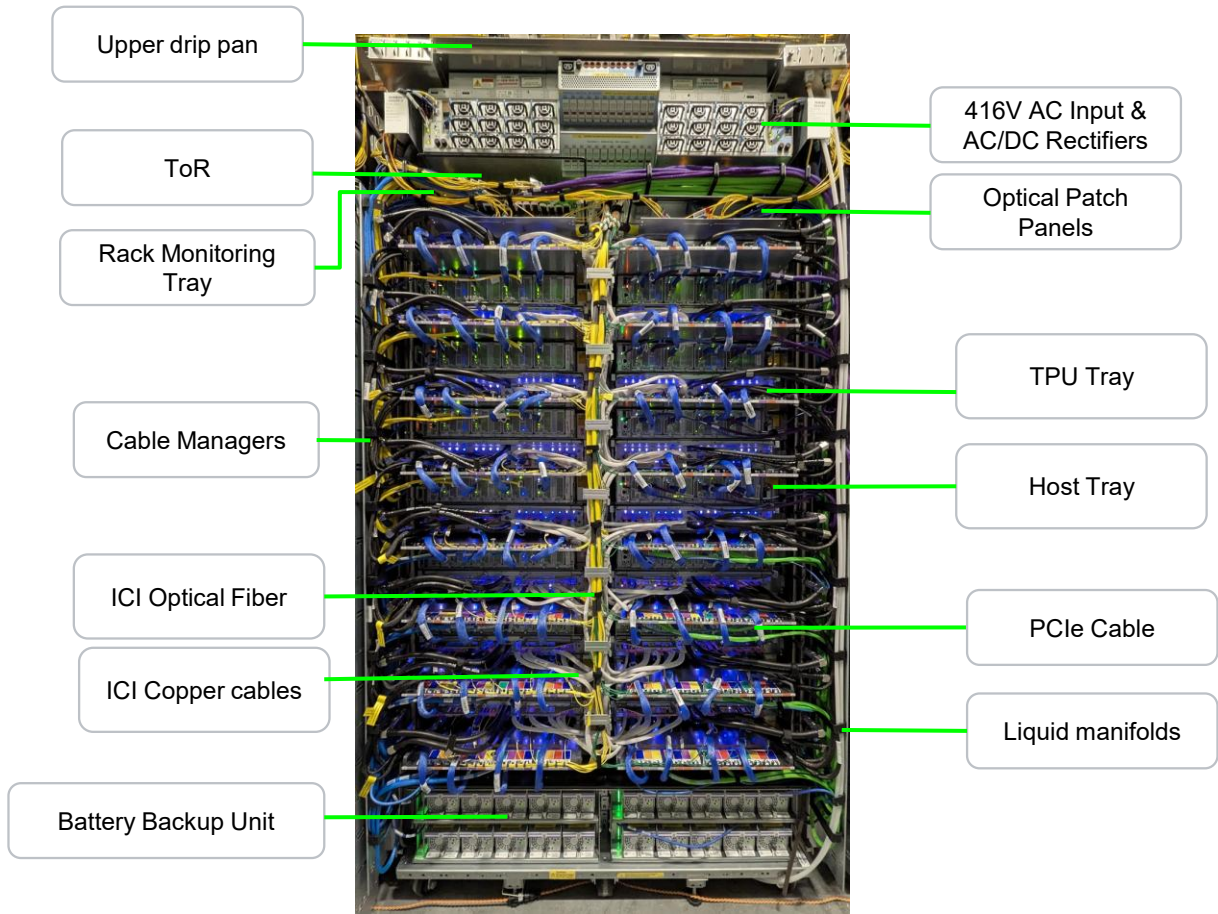
Ironwood Rack

Rack is the Building Block

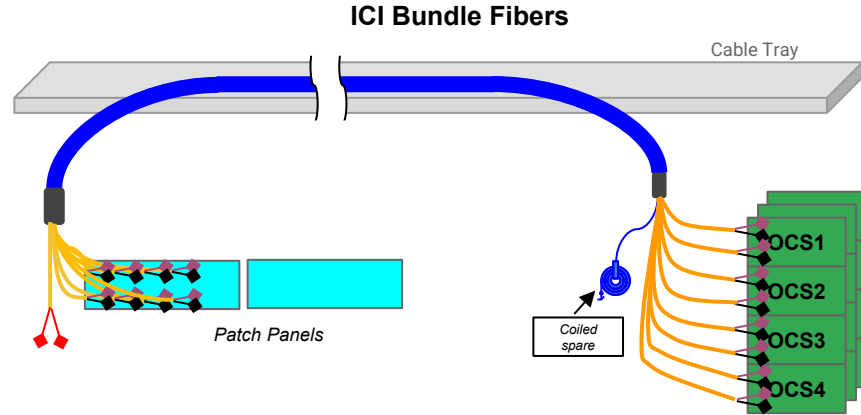
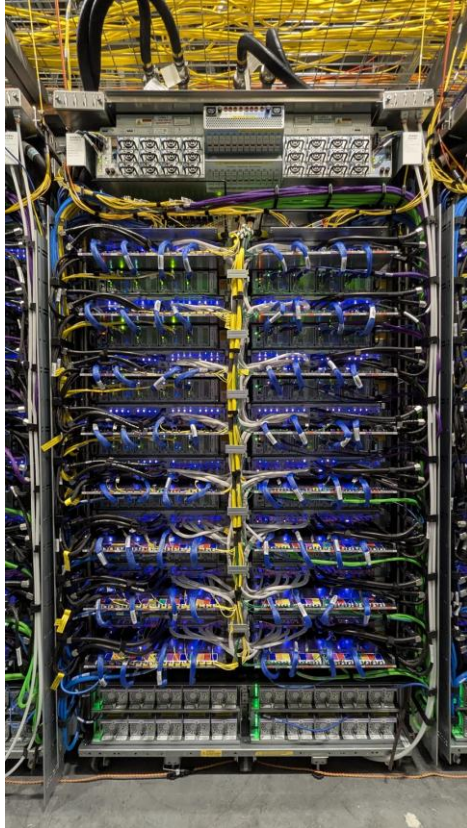


- 4x4x4 building block over ICI that fits in a single Google TPU rack(wider than a typical rack)
- 6 faces on the cube, each with 16 links
- ICI connections: 96 Optical fiber, 80 Copper cables & 64 pcb traces

Rack Subsystems

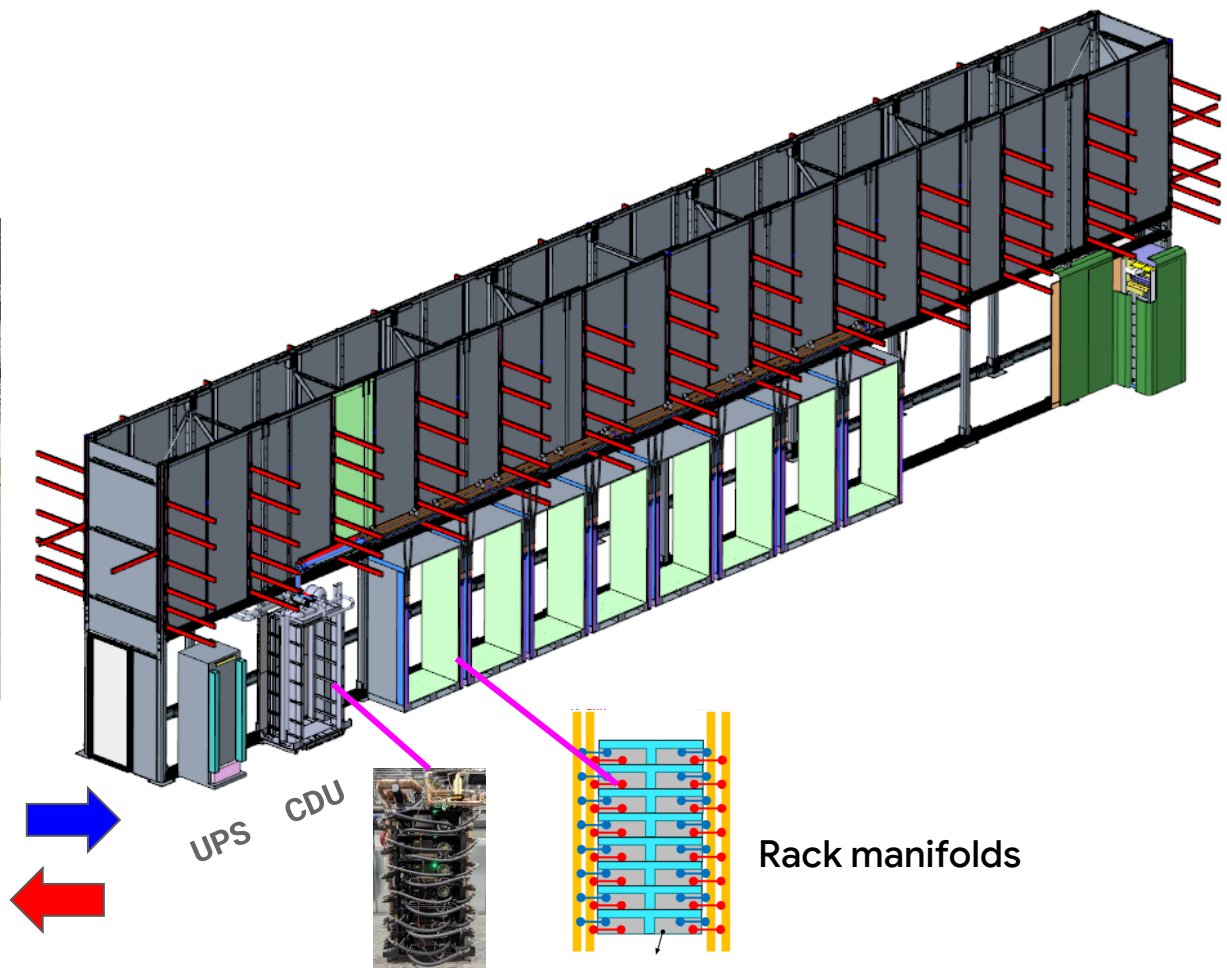
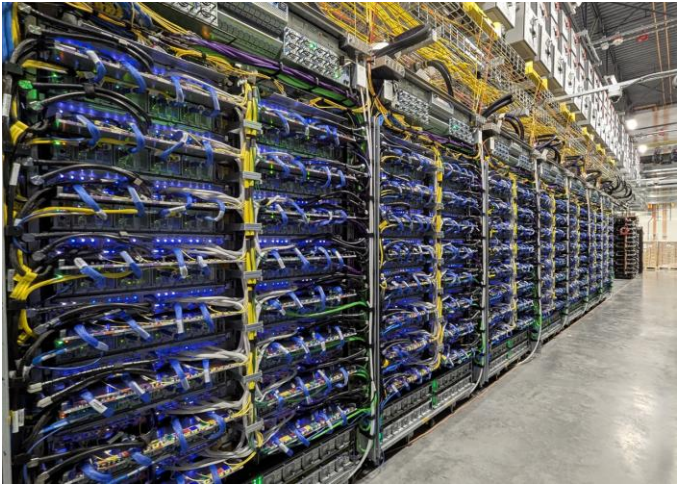


ICI Fiber Bundles



Pre-deployed Fiber for ICI Fiber Bundles

Cooling Infrastructure & DC Layout



Facility PCW

UPS

CDU

Rack manifolds

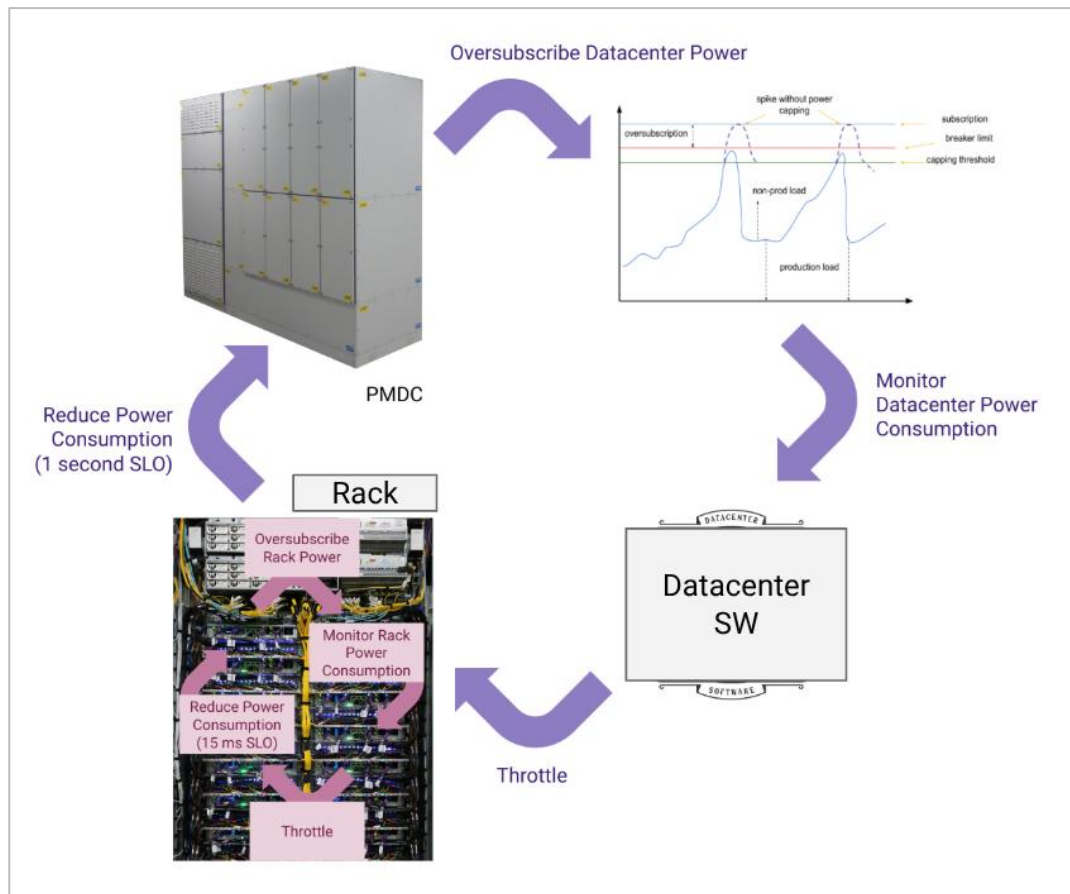
Power Management

TPU SoC Power Capping

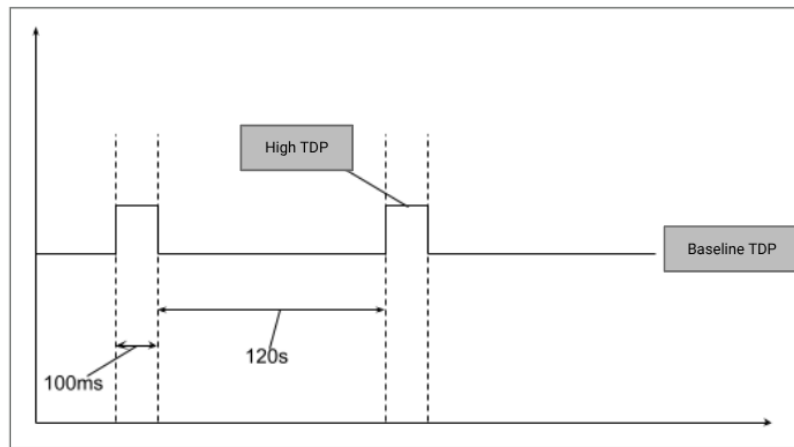
- TPU Power Capping prevents ML jobs from exceeding datacenter power provisioning thresholds

Rack Capping

- SLO of < 15 milliseconds.
- Support baseline TDP and High TDP modes to unlock the peak TPU Perf



Rack Power Capability Boundary



- Continuous Baseline TDP Mode (time scale > 1s)
- High TDP Mode(100ms)
- Rack power throttling mode lasts for 120s when it is activated.

Challenges & Opportunities

Challenges

- End of Dennard Scaling has led to an increase in power(gen over gen)
- Front Panel Cabling can be congested as our I/O requirements increase
- High Power Swings under Synchronous Training

Opportunities

- Collaborate on industry-wide ML rack standards
- Develop shared infrastructure for future ML racks
- Share best practices for ML rack deployment