

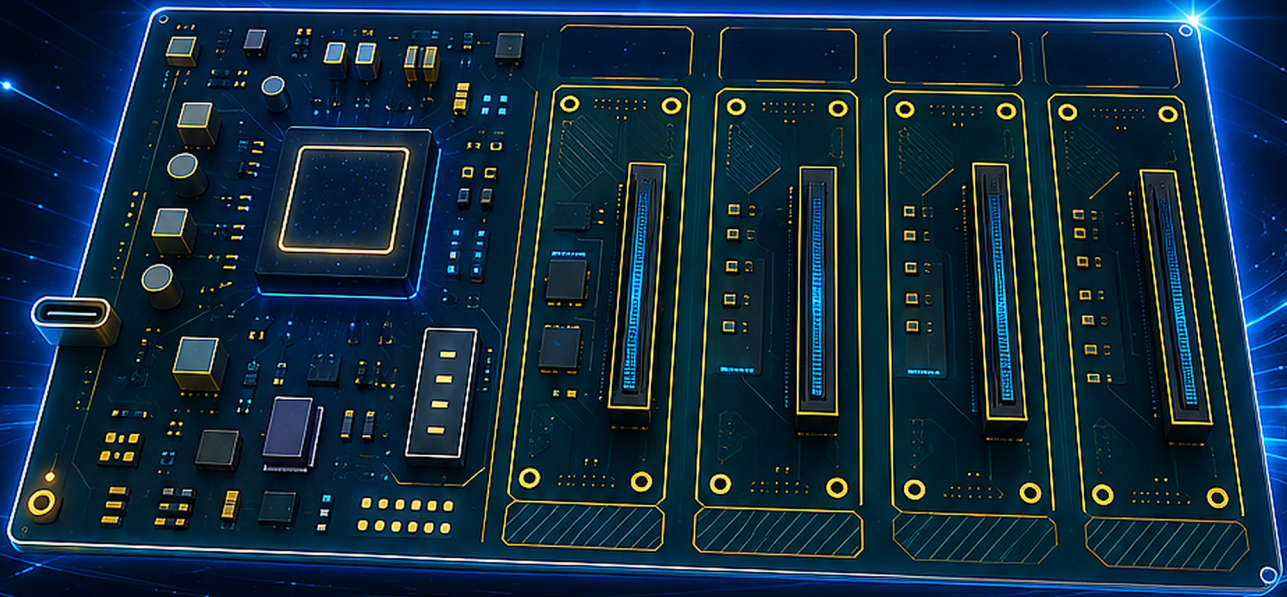


PRISME

PRISME Q8

Data should not hold your GPU back.

Move less. Reduce CPU overhead. Feed the GPU earlier.



OPTIMIZE YOUR GPU. WITHOUT REPLACING YOUR PC.

PRISME Q8 is designed to move less data work through the CPU and prepare selected data earlier for GPU compute.

6.4 GB/s

INTERNAL DESIGN RATE

256 bytes x 25 MHz. Host/DMA throughput is measured in the PoC.

GPU COMPANION

FIRST MODULE

Prepares and verifies data before GPU compute. Effect is measured per workload.

ONE BASE

FOUR MODULE PATHS

Shared control model. Expand when the use case requires it.

OPTICAL

DATA PROCESSING

Designed for fewer unnecessary data movements. Benefits are measured per workload.

BUILT FOR:

GAMING

Measure FPS, 1% lows and frame-time.

AI & GENERATIVE

Measure inference, throughput and CPU overhead.

LOW LATENCY

Measure response, jitter and deterministic timing.

VIDEO & CREATIVE

Measure workflow time and J/GB.

DESIGN TARGETS. FINAL MINIMUM GUARANTEES FOLLOW POC AND INDEPENDENT VALIDATION.

WHY CPU + GPU

STILL WASTE TIME

Too much unnecessary work. Too late.

Bottlenecks, waiting and uneven performance can appear between the compute stages.



1 CONVENTIONAL PATH



1 Too many copies

Data may be copied multiple times between CPU and GPU.



2 Unnecessary CPU work

CPU time may be spent on movement, formatting and coordination.



3 GPU waits for data

The GPU may stand idle while data is prepared.



4 Uneven frame-time

Spikes can affect 1% lows and responsiveness.

Less waste. More useful work.



2 PRISME PATH



1 Earlier verification

Data is validated and prepared before the next heavy stage.



2 Less data movement

The target is fewer selected copies and fewer hops.



3 More useful work

CPU time is freed for tasks that actually require the CPU.



4 Closer to the GPU

A shorter logical path to the next compute stage.

PRISME Q8 IS DESIGNED TO REDUCE UNNECESSARY INTERMEDIATE STEPS.

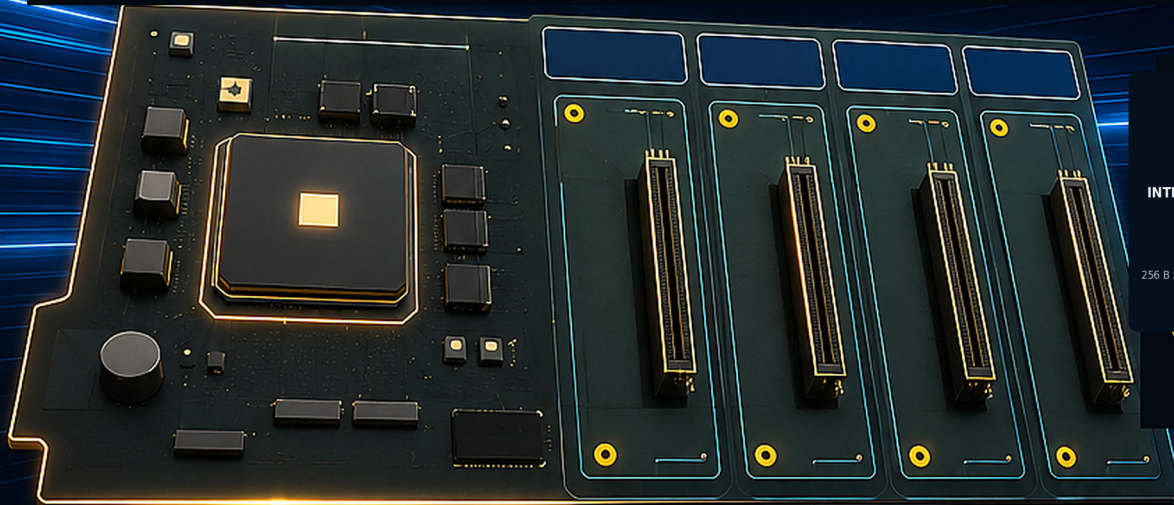
Measure throughput, latency, CPU/GPU load and frame-time on the same system and the same workload.

PRISME Q8

GPU COMPANION

FOR MAINSTREAM PCs

Built to reduce CPU overhead and feed the GPU earlier.



6.4 GB/s

INTERNAL ARCHITECTURE
BASELINE

256 B x 25 MHz. Host/DMA is measured
in the PoC.



GB/s throughput

More data per second



CPU load

Lower CPU overhead



GPU utilization

Higher accelerator use



1% low / frame-time

Stability where it matters



Watt

Power per workload



J/GB

Better efficiency

ivitet

POC BENCHMARK: SAME SYSTEM. SAME WORKLOAD. WITH / WITHOUT PRISME.

GB/s

throughput

CPU %

overhead

GPU %

utilization

1% LOW

frame-time

W

power

J/GB

efficiency

No percentage is published as a guarantee before independent validation.

TARGET: HIGHER GPU
UTILIZATION

The GPU receives data earlier and more consistently; the effect is measured in real-world workloads.

TARGET: LOWER CPU OVERHEAD

Q8 handles relevant preparation and verification where the architecture benefits from it.

TARGET: BETTER FRAME-TIME

Measure frame-time, 1% lows and spikes - not just average FPS.

TARGET: BETTER EFFICIENCY

Measure watts and J/GB together with throughput and actual workload results.

MEASURABLE PERFORMANCE. GUARANTEES AFTER INDEPENDENT VALIDATION.

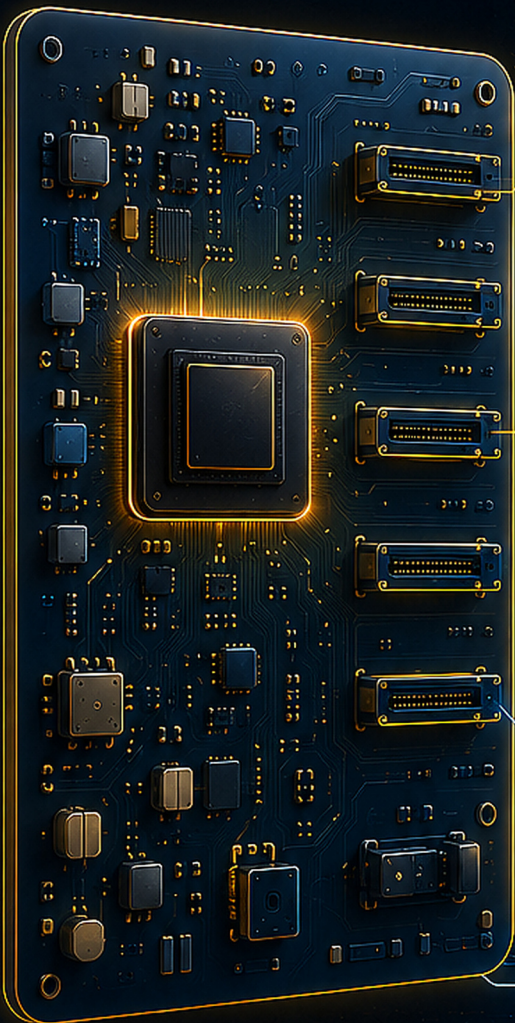
Measure -> validate -> document -> define conservative minimums for the specific configuration.

LESS WASTE. MORE USEFUL WORK.

ONE BASE - FOUR PATHS FORWARD

PRISME Q8 is a modular platform.

Buy one platform. Expand when the use case creates value.



01

GPU COMPANION

FIRST COMMERCIAL MODULE



Prepares and verifies data before GPU compute via PCIe. Designed for lower CPU overhead and more stable data feeding.

02

FACTORY PROGRAMMER



Programming, read-back/verify and production logging. Adapters for ESP32, Arduino-class, Raspberry Pi and other targets.

03

LOCAL LATENCY 2-300 m



Optical delay/buffer. Approx. 9.8 ns-1.47 us fiber-only. Electronics latency is measured separately.

04

DATACENTER



Standard PCIe server card + 100G optical. External 8-10 km latency appliance. Same PRISME protocol.

WHY IT SCALES

SHARED PLATFORM

Same software and control model.

MODULAR INTEGRATION

New functions are added as modules.

BACKWARD COMPATIBLE

Backward compatibility is a design principle.

START WITH THE NEED

Expand when the use case creates value.

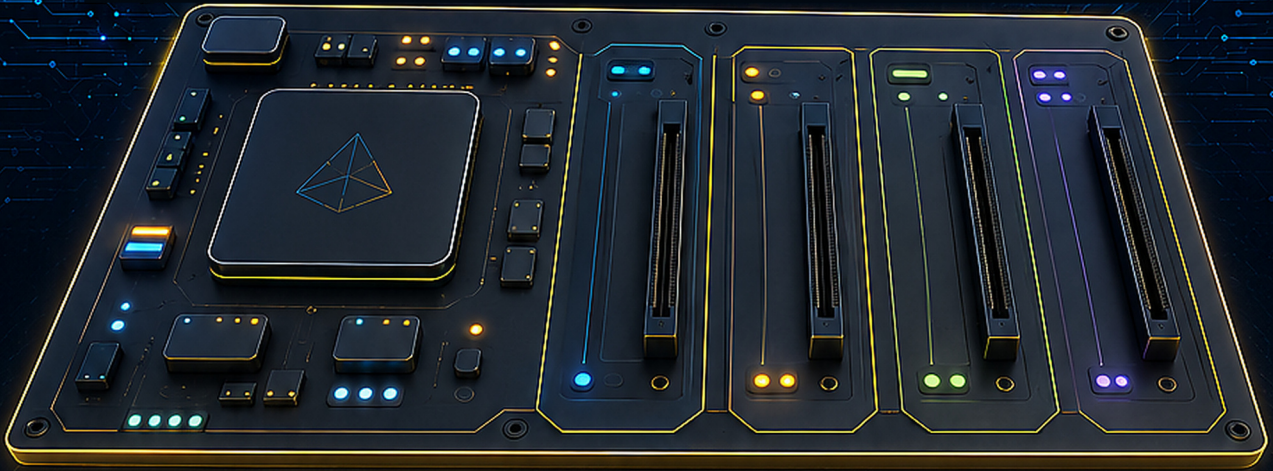
FUTURE ACCELERATOR READY - GPU · LPU-CLASS · AI · NEXT-GENERATION COMPUTE.



PRISME Q8

LAUNCH PARTNER / LOI

Be among the first to evaluate PRISME Q8 and register volume interest.



SCAN FOR LOI / SALES PAGE

EARLY ACCESS

Access pilot participation, technical material and the validation process.

VOLUME INTEREST

Register expected volume for the first production batch.

INPUT ON PRIORITIES

Feedback can influence module prioritization and the roadmap.

DOCUMENTED VALIDATION

Receive test results and validation dossiers when completed.

MINIMUM INITIAL ORDER: 50 UNITS

PRISME Q8 - MODULAR COMPUTE FOR PC, INDUSTRY AND DATACENTER.

GitHub

github.com/janusSG/PRISME

Demo

janusSg.github.io/PRISME/demo/prisme.html

Contact

Via the LOI page / QR code

LAUNCH PARTNER MATERIAL - PRE-PRODUCTION. PERFORMANCE GUARANTEES FOLLOW VALIDATION.