

GPU sharing with HAMi



Founds · Maintains · Drives adoption



GPUs & accelerators



Capabilities vary by device.

HAMi in the cloud native ecosystem



Contributors

723

Organizations

1075

Countries and regions

44

Incubating since July 2026



CLOUD NATIVE
COMPUTING FOUNDATION



KAI Scheduler



HAMi-core

Per-container GPU memory limits

Schedulers



VOLCANO



koordinator



Kueue

Cloud & AI platforms

UCLLOUD



DaoCloud



金山云

HAMi in production



招商銀行
CHINA MERCHANTS BANK

+46.7%

Single-card inference throughput
ResNet 50 · batch 32

SNOW

30% fewer GPU hours

Driving simulation



50% fewer GPUs

Train-to-inference pipelines



90% of GPU infrastructure

Managed by HAMi

HAMi adopters



Dynamia's engineering contributions



MetaX integration

Prefill / decode compatibility

Under review

Lupine

vLLM compatibility

Per-client GPU metrics

Merged upstream

Later today · 8 Sep (UTC+8)

09:59–10:04

PD Disaggregation vLLM Deployment
on Alternative AI Accelerators Using llm-d

11:14–11:19

Project Lightning Talk: From Static Slices
to Elastic GPUs: Dynamic MIG with HAMi

14:30–15:00

How Intsig Serves Billions of Document Scans:
GPU Virtualization at Scale with HAMi

HAMi community

WeChat assistant

