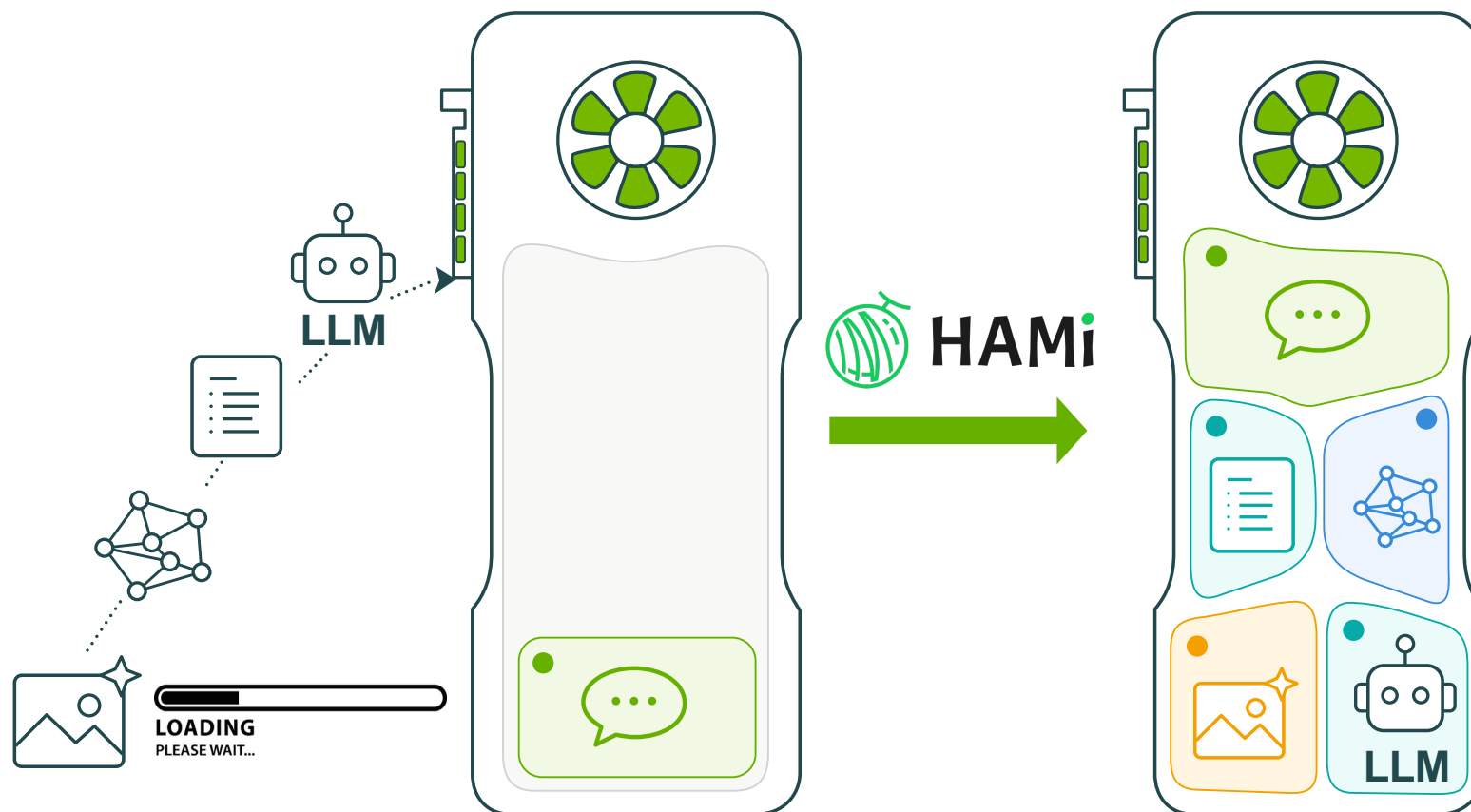


GPU sharing with HAMi



GPUs & accelerators

AMD

Ascend

aws

壁仞科技
BIREN TECHNOLOGY

Cambricon
寒武纪

Enflame
燧原科技

HYGON
中科海光

天数智芯
iluvatar CoreX

昆仑芯
KUNLUNXIN

META
沐曦

摩尔线程
MOORE THREADS

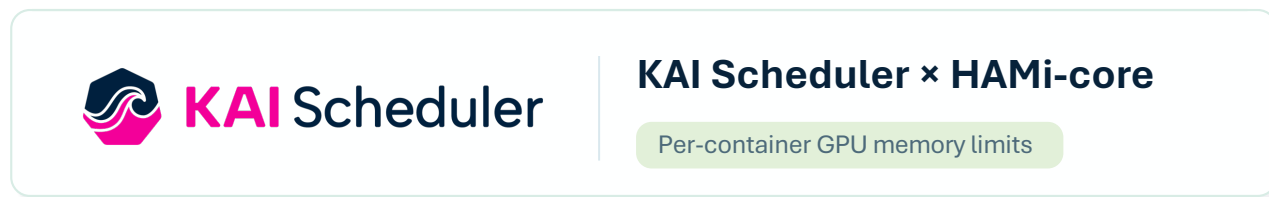
NVIDIA

瀚博半导体
POWERED BY VASTSTREAM

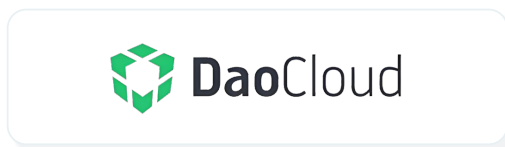
Capabilities vary by device.

HAMi in the cloud native ecosystem

Schedulers



Cloud & AI Platforms



723

Contributors

1075

Organizations

44

Countries & Regions



Incubating since July 2026

HAMi in production



+46.7%

Single-card inference throughput

ResNet50 · batch 32



30%

fewer GPU hours

Driving simulation

SNOW

50%

fewer GPUs

Train-to-inference pipelines

PREP

90%

of GPU infrastructure

Managed by HAMi



Dynamia's engineering contributions



MetaX integration

Prefill / decode compatibility

Under review

Lupine

vLLM compatibility

Per-client GPU metrics

Merged upstream



Founded by HAMi initiators
Maintains & drives
HAMi adoption



CNCF Incubating Project

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

Booth area: Grand Ballroom I, Project Pavilion, **T-1**

HAMi community

WeChat assistant

