

Scaling Distributed Inference with llm-d

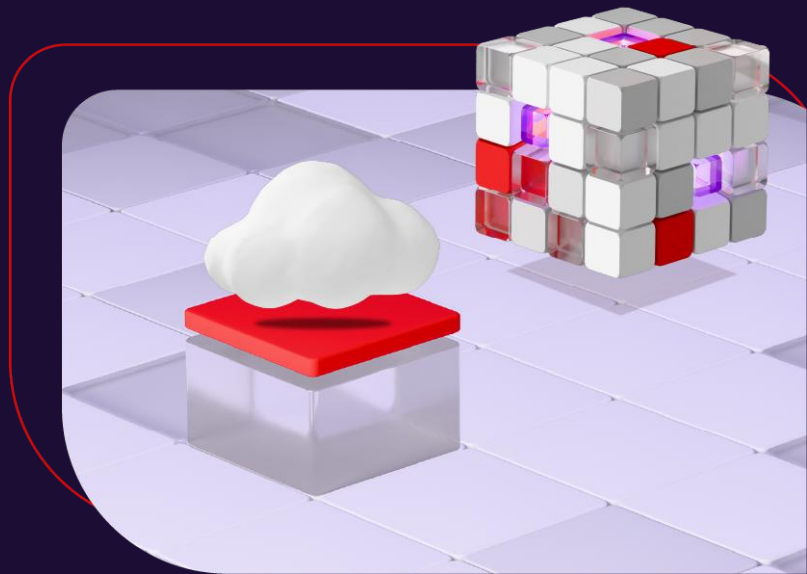
vLLM Meetup Bangkok, 6 Sep 2026

Jing Wen Ng

Specialist Solution Architect, AI, Asia Pacific @ Red Hat



The Problem



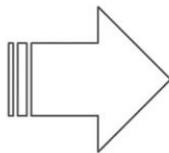
Why Distributed Inference?

Modern HTTP requests:

Fast, uniform, cheap

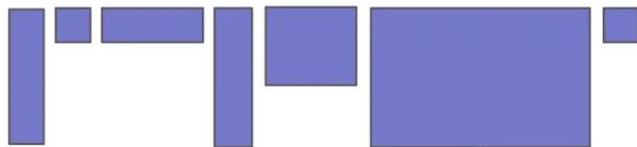


3-4 orders of magnitude more QPS



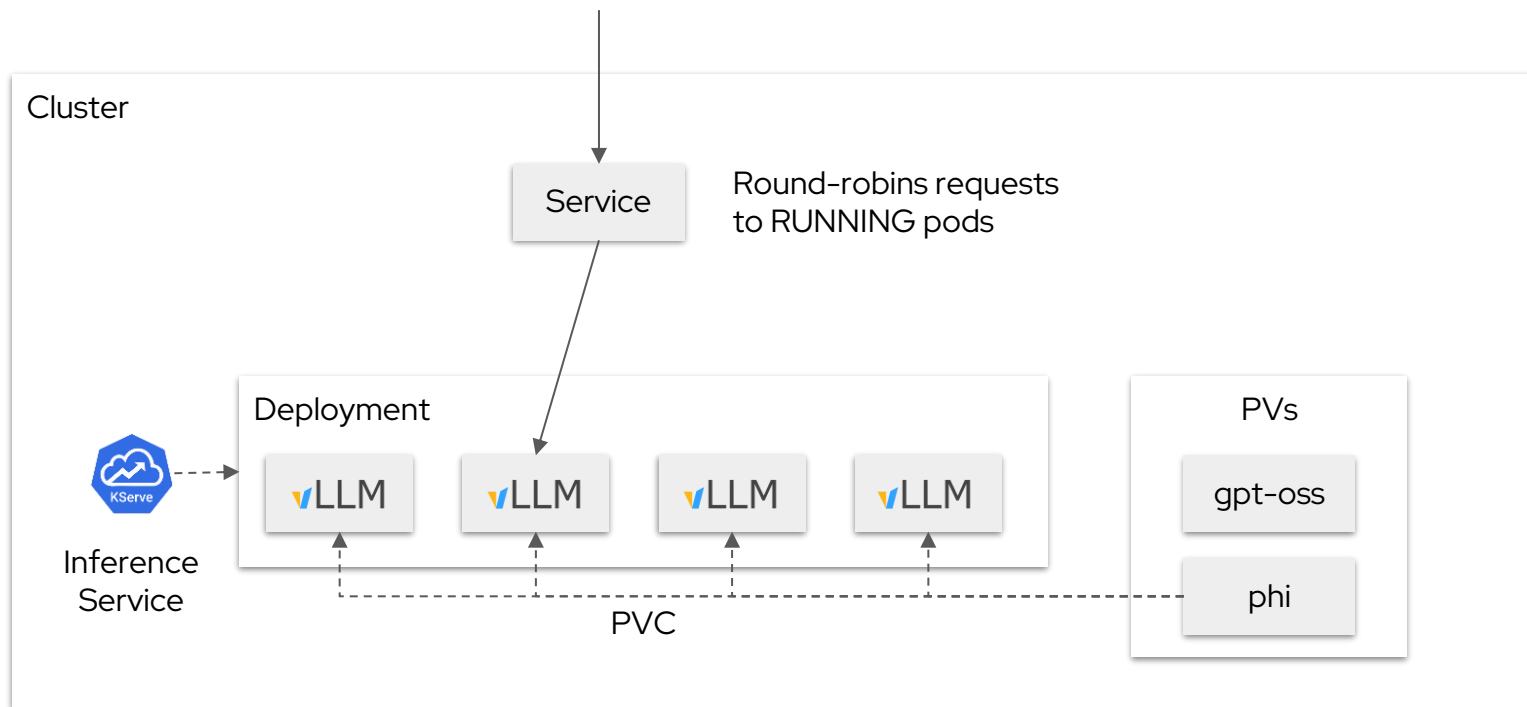
LLM requests:

Slow, non-uniform, really expensive

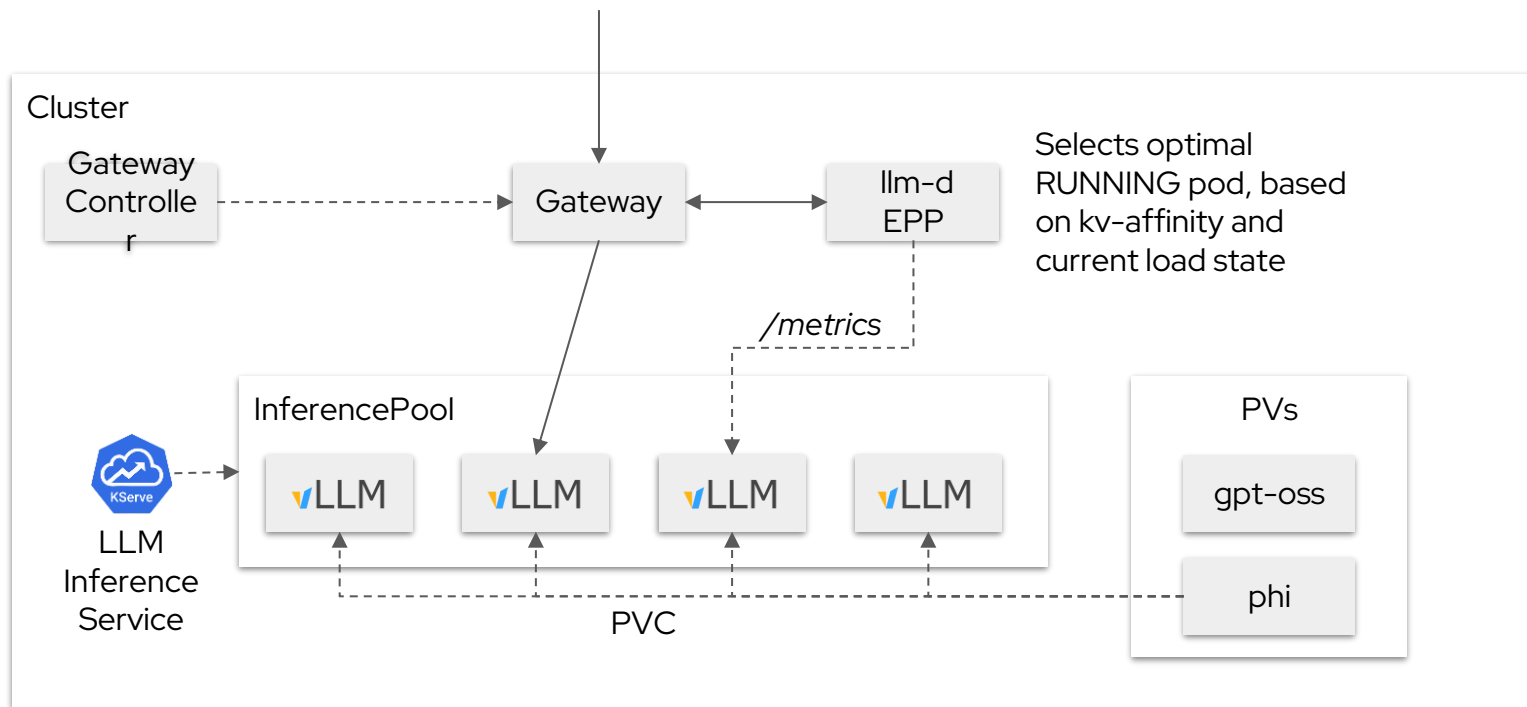


6-9 OoM more expensive per request

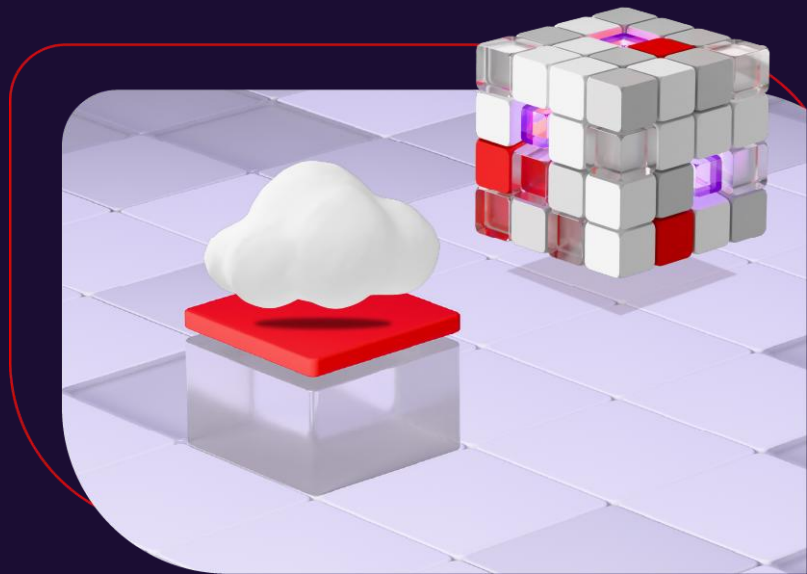
Baseline Deployment Architecture



RHOAI llm-d Deployment Architecture



What is llm-d?



What is llm-d?

An open-source, Kubernetes-native framework for distributed LLM inference

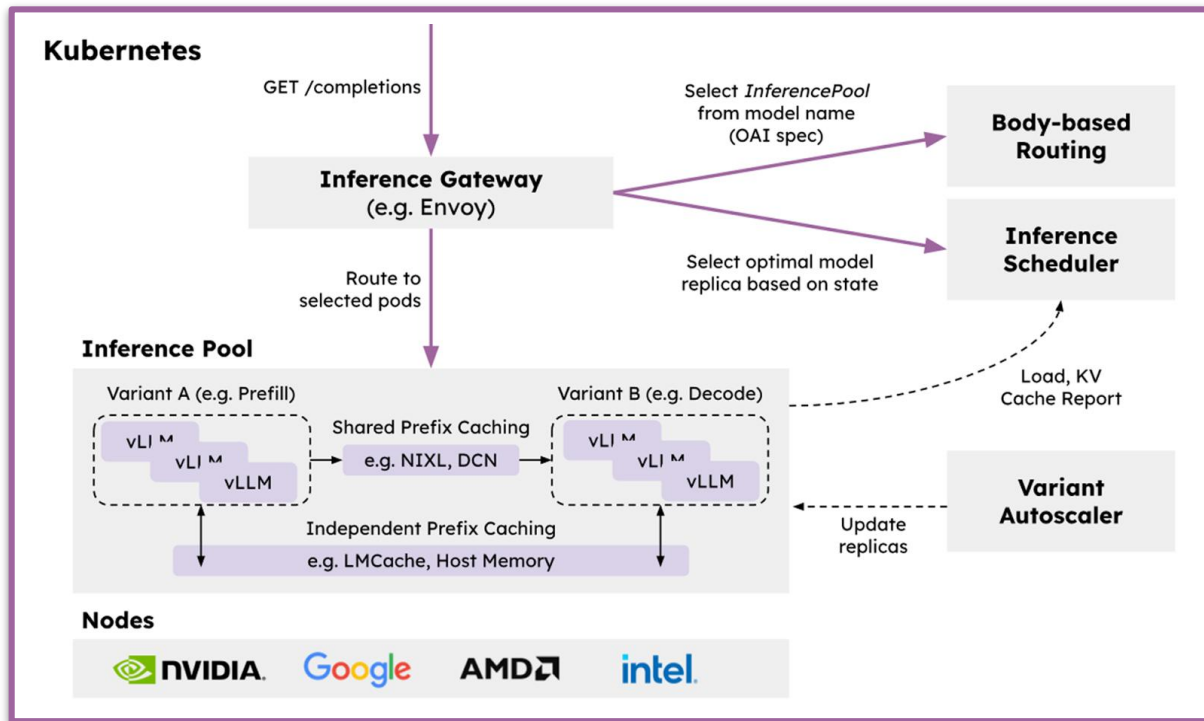


An open-source framework for distributed large language model (LLM) inference that **runs natively on Kubernetes.**

- ▶ Joint open source initiative by Red Hat, Google, NVIDIA, Hugging Face, and many more organizations
- ▶ Optimized for Kubernetes based distributed platforms (OpenShift)
- ▶ Designed to optimize scale, latency, and flexibility for AI workloads
- ▶ Built to extend and interoperate with vLLM
- ▶ Maximizes performance across multi-tenant and multi-model workloads with intelligent scheduling and resource management



llm-d Architecture



- **Operationalizability:** modular and resilient architecture with native integration into Kubernetes via Inference Gateway API
- **Flexibility:** cross-platform (active work to support NVIDIA, Google TPU, AMD, and Intel), with extensible implementations of key composable layers of the stack
- **Performance:** leverage distributed optimizations like **disaggregation** and **prefix-aware routing** to achieve the highest efficiency for available infra

Key Features (“Well-Lit” Paths)

sett

ing

s_s

ugg

Intelligent Inference Scheduling

KV-cache aware, load-aware routing

call

_sp

lit

Prefill-Decode Disaggregation

Decoupled execution for higher throughput

traf

fic

Flow Control

Multi-tenancy & fairness-scheduling

me

mo

ry

KV Cache Management

CPU-offloading & Storage-offloading

lan

Multi-Node Serving

“Wide Expert Parallel” serving architecture

tre

ndi

ng_

up

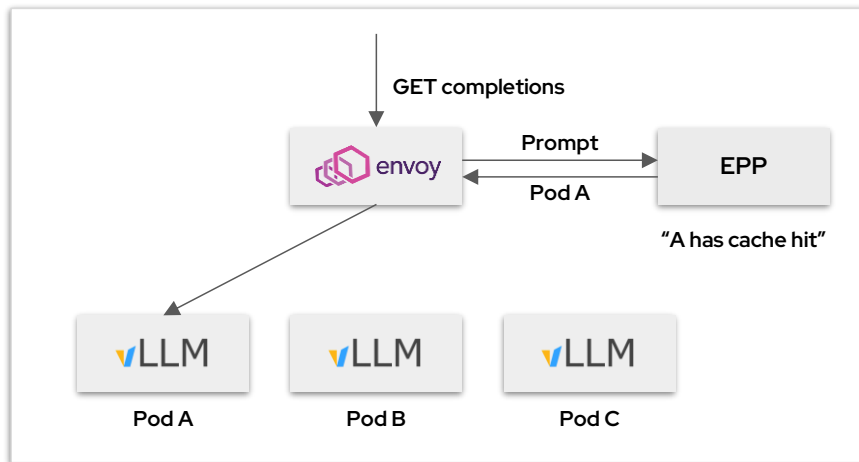
Workload Variant Autoscaling

Dynamic and highly responsive auto-scaling

Intelligent Inference Scheduling

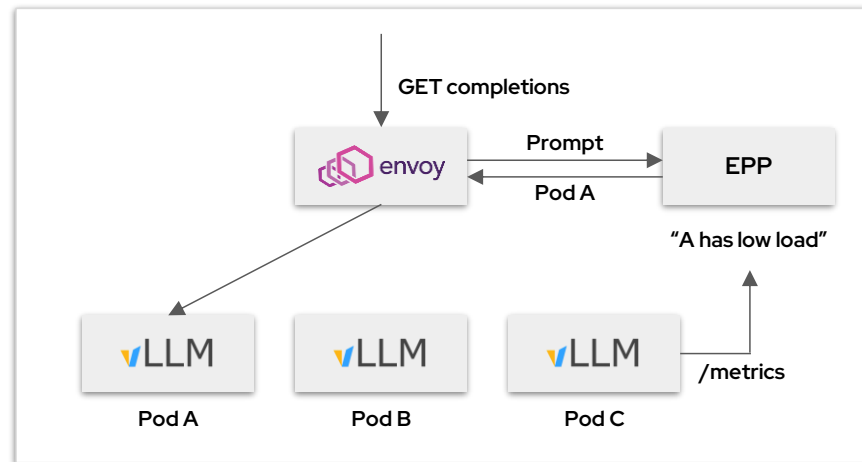
vLLM-aware load-balancing enables smarter request routing that improve SLOs

Prefix-Aware Routing



Dramatically increase prefix-cache hit rate

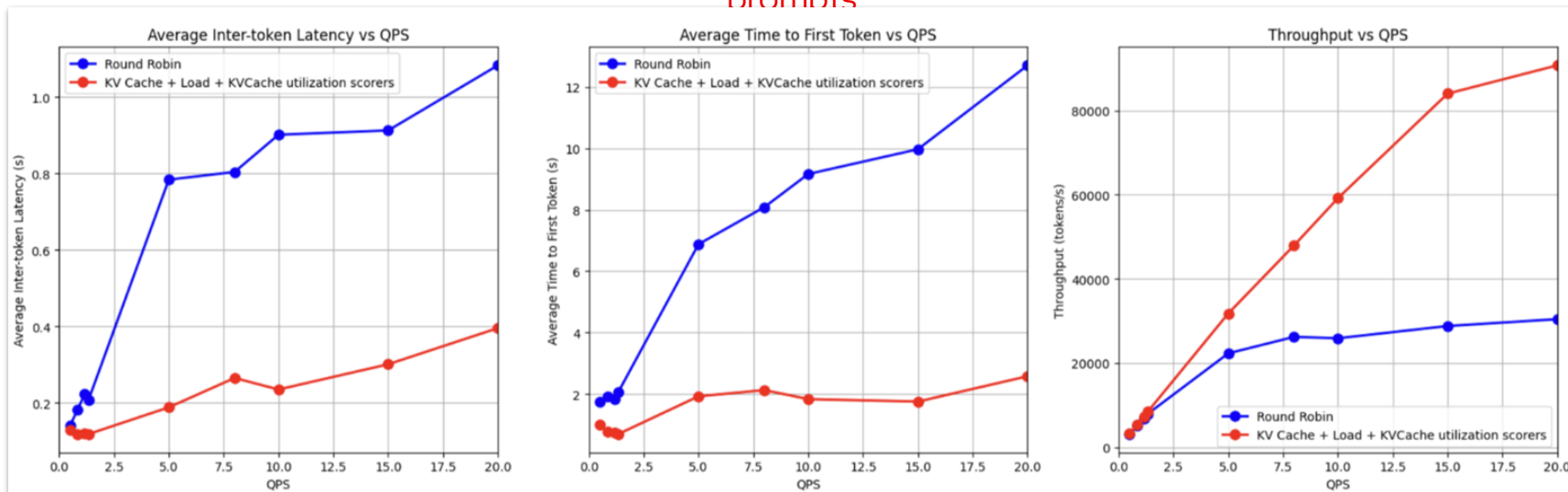
Load-Aware Routing



Load-balancing based on actual replica state

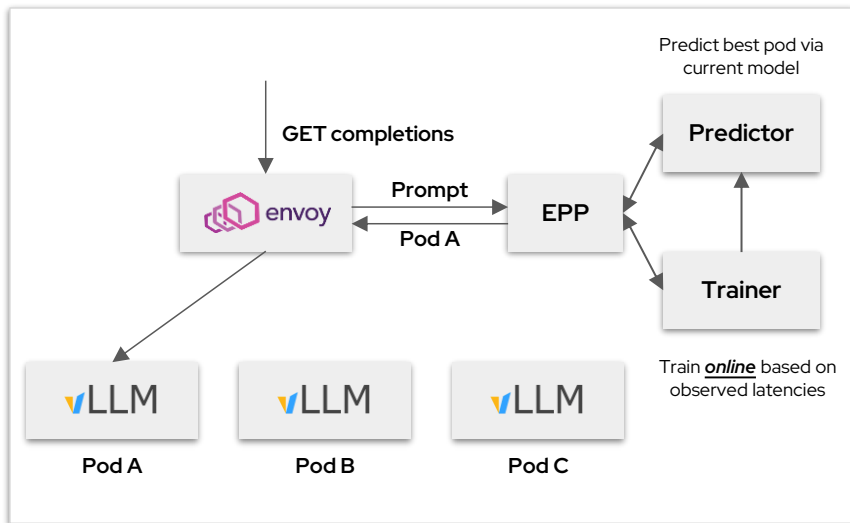
Intelligent Inference Scheduling

Inference scheduling is a no-brainer optimization which can have huge impacts on repeated prompts



Predicted Latency Scheduling

We are currently improving the Intelligent Inference Scheduling with predicted latency



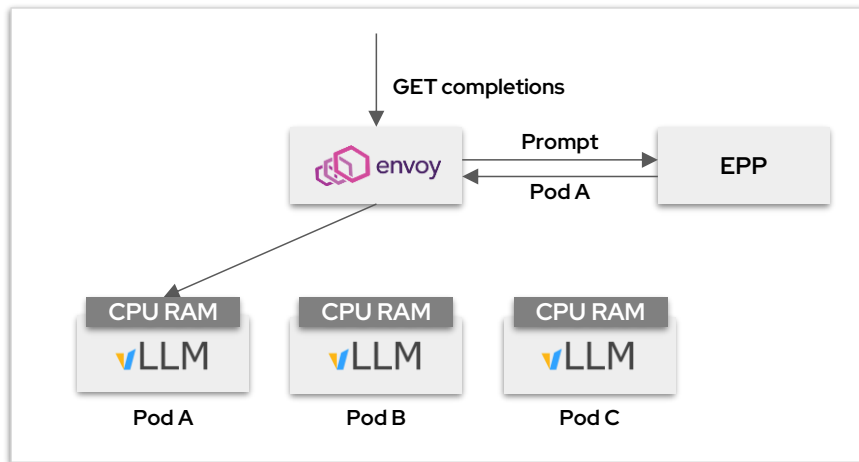
Scenario	Success Rate	E2E p50	E2E p95	TTFT p50	TTFT p95	TPOT p50	TPOT p99
No Gateway	~99.8%	15.98s	38.85s	4.47s	24.04s	35ms	93ms
Load + Prefix (111)	~99.9%	16.42s	35.06s	2.86s	18.06s	39ms	103ms
Load + Prefix (322)	100%	13.42s	26.55s	3.38s	16.78s	28ms	63ms
Latency Prediction	~99.9%	9.06s	22.57s	0.97s	11.34s	22ms	53ms

<https://llm-d.ai/blog/predicted-latency-based-scheduling-for-llms>

KV Cache Management

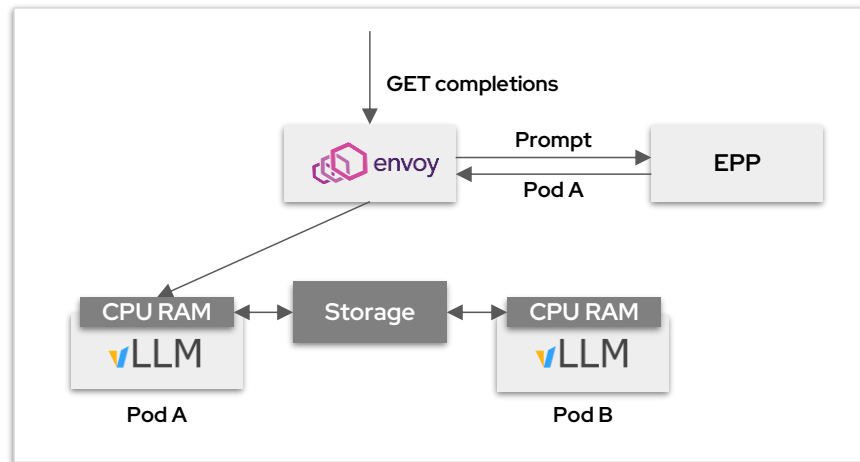
Leverage all system resources to maximize prefix cache hit rate within the cluster

CPU KV Cache Offloading



Offload KV caches to local memory with global index

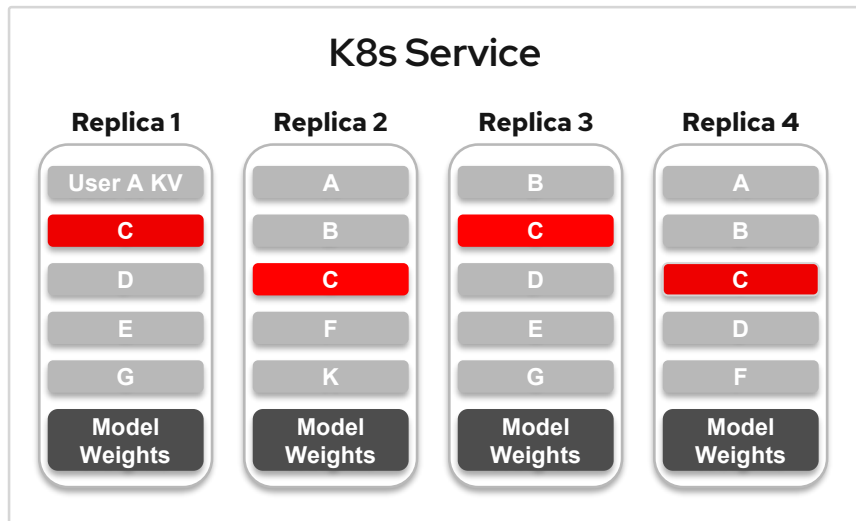
Storage KV Cache Offloading



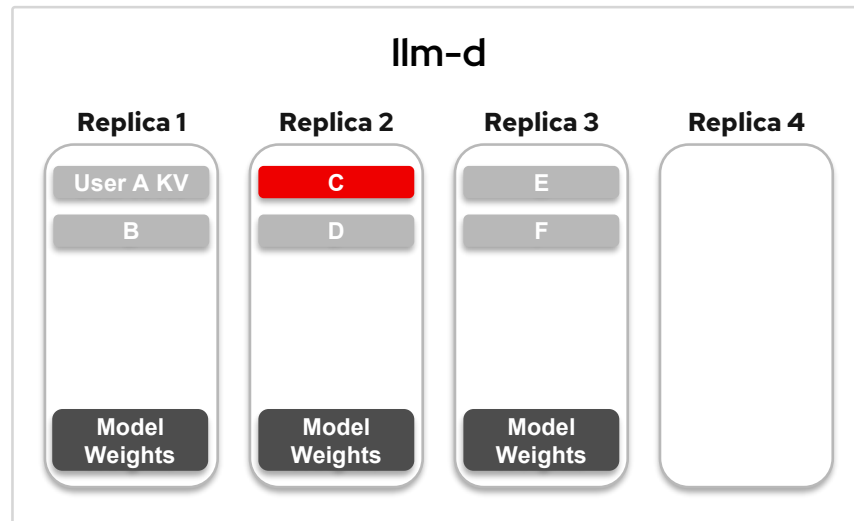
Offload KV caches to global shared remote storage

Example: Impact on KV Cache Re-Use

Leverage all system resources to maximize prefix cache hit rate within the cluster



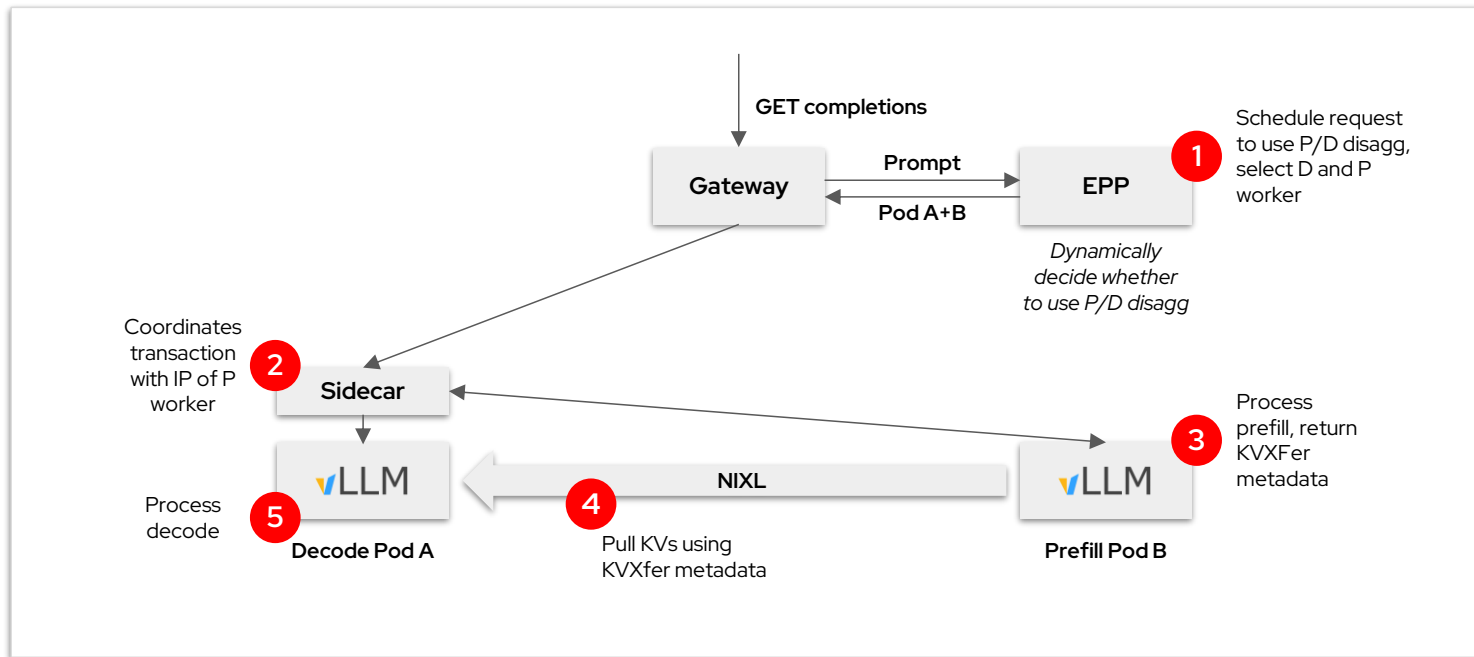
User C's requests are round-robin
across all replicas



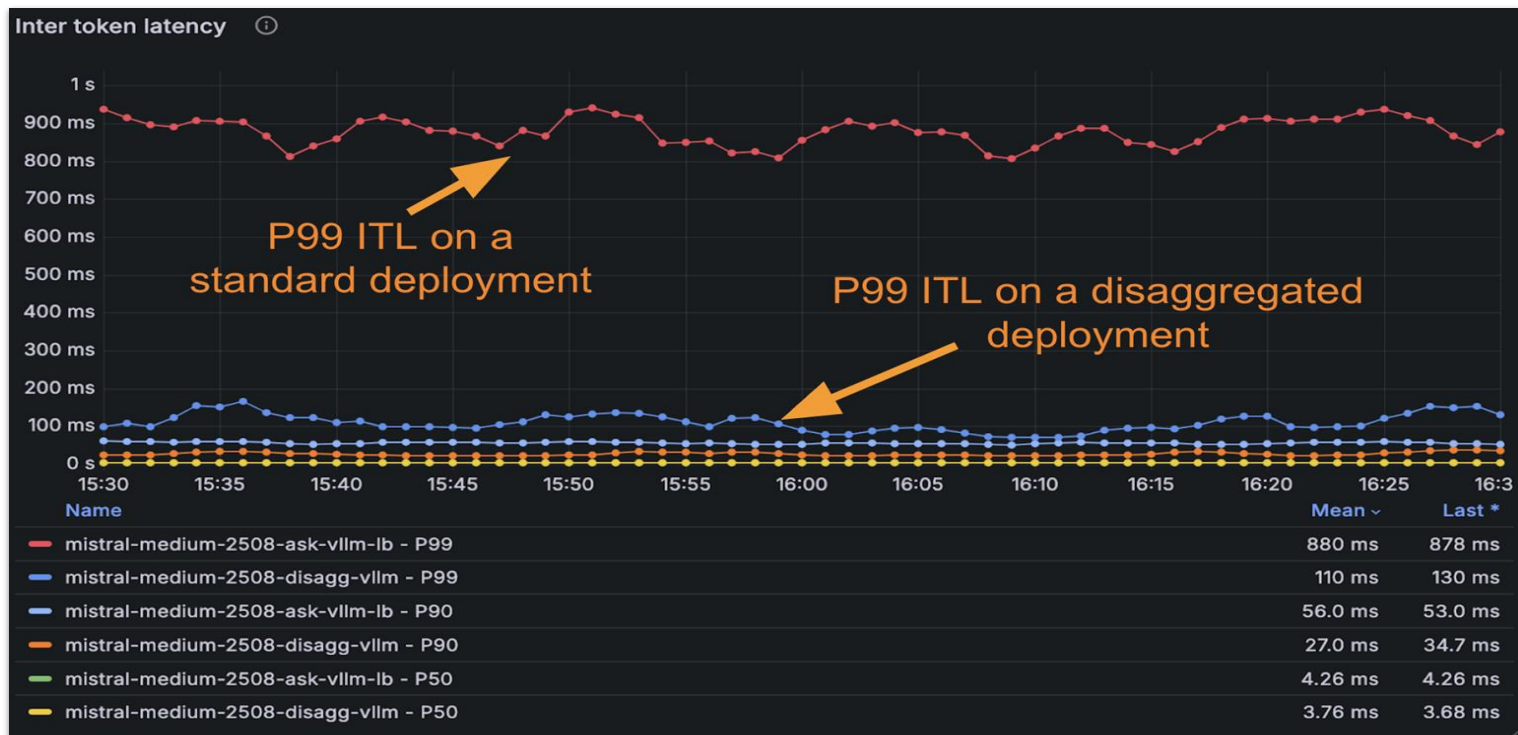
User C's requests routed to pods with cached
tokens, **skipping redundant prefix**

P/D Disaggregation

llm-d composes vLLM and Gateway to enable P/D disaggregation

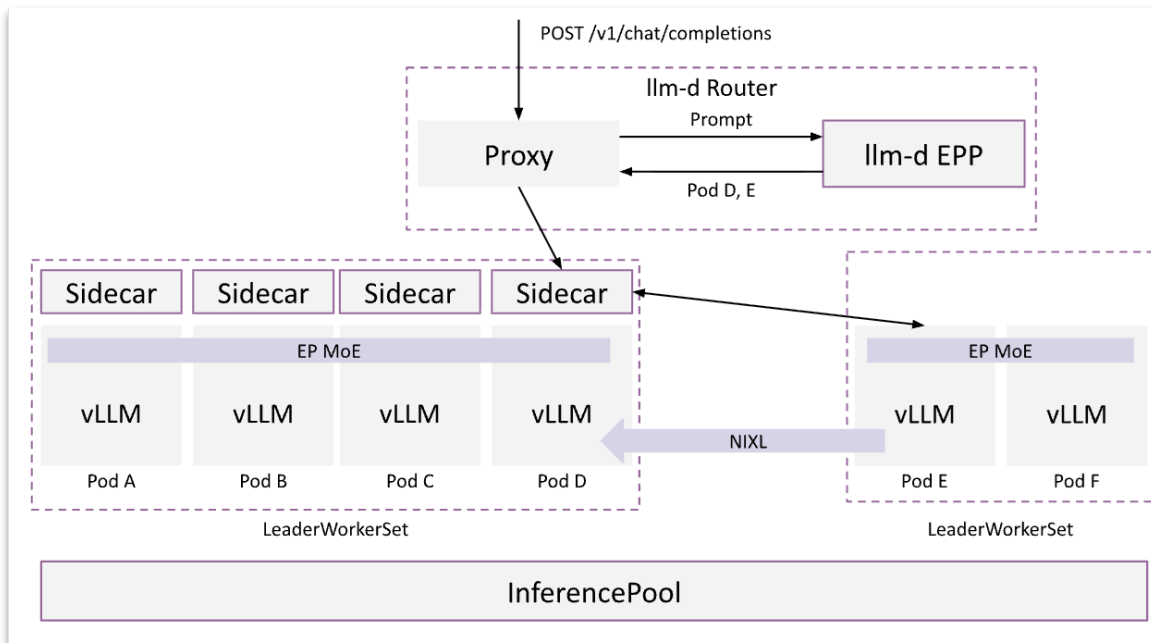


P/D Results - Inter-Token Latency (ITL) smoothing

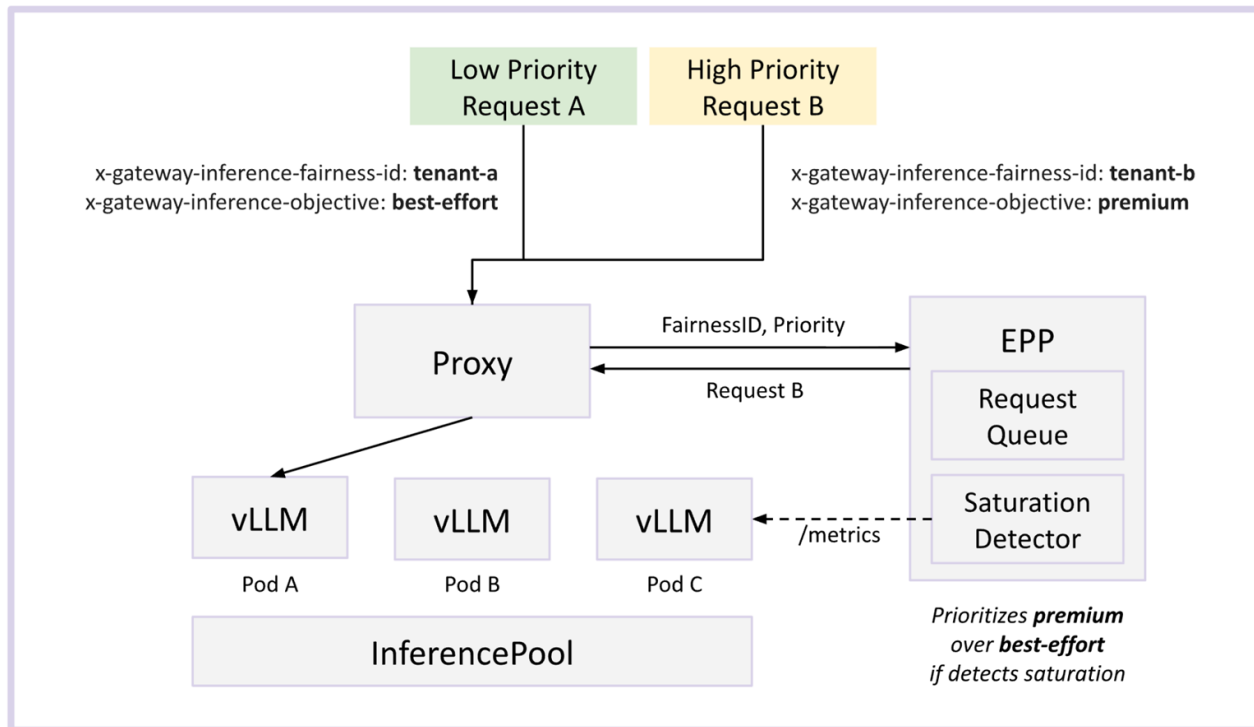


Multi-Node Wide Expert Parallelism

llm-d's K8s-native design composes the DP/EP implementation with the rest of the llm-d system

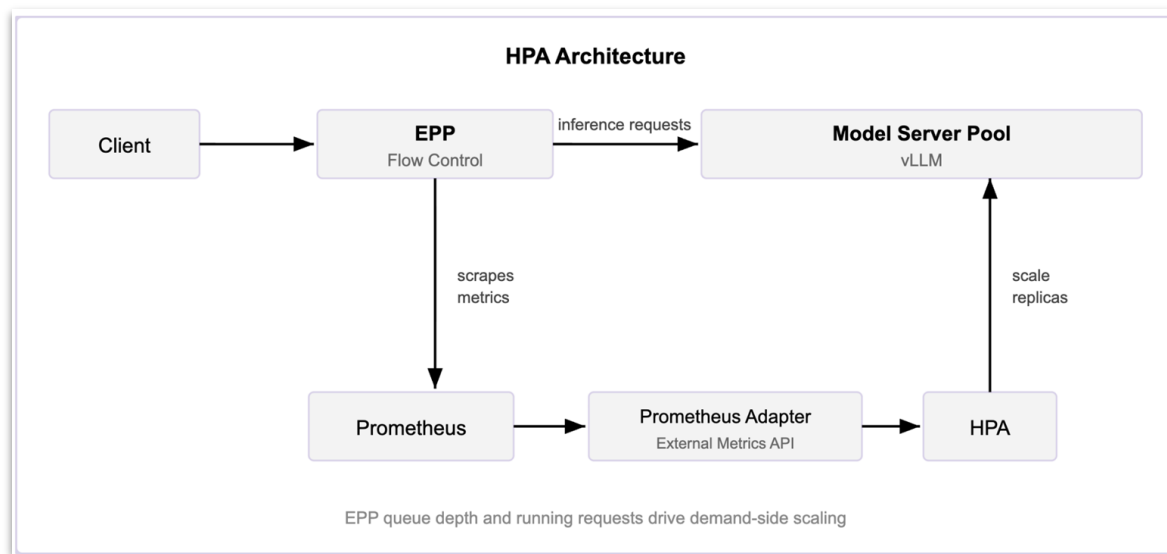


Flow Control for Multi-Tenant Use Cases



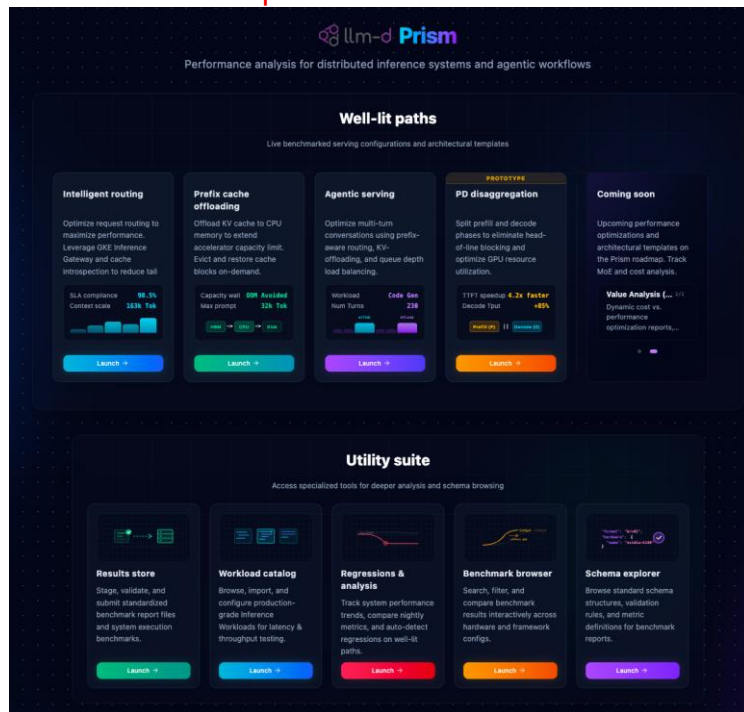
Autoscaling

Flow control multiplexing of different request classes onto the same model deployment

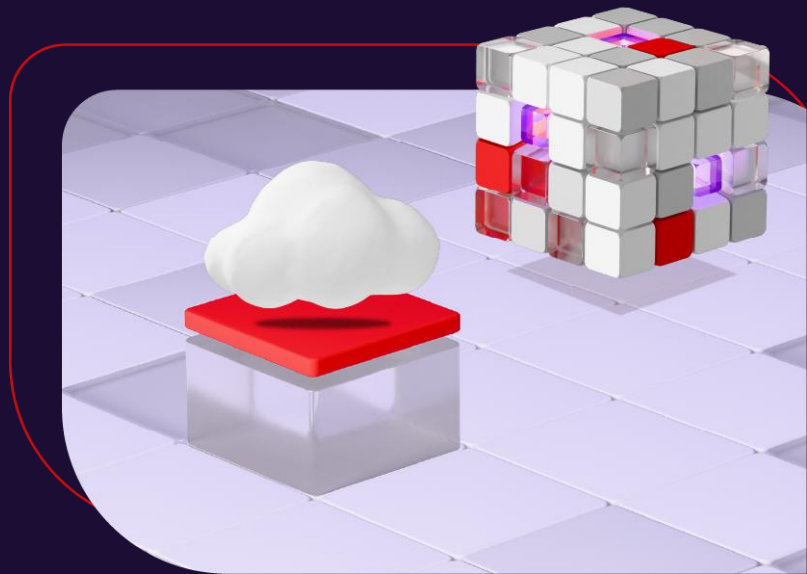


Well Lit Paths on Prism.llm-d.ai

Visualize and compare metrics across benchmarks



Essential Elements of Enterprise GenAI Stack



Enterprise GenAI Inference Platform

Operationalize Model as a Service across the hybrid cloud

AI Gateway

API Management, Token Tracking,
Authentication for any model



Distributed Inference



Nemo



GPT OSS



Gemma



Llama



Granite



Mistral



DeepSeek



Qwen

Validated, Optimized Models &
Techniques



Inference Server



NVIDIA GPUs



AMD GPUs & CPUs



TPUs



AWS Inferentia



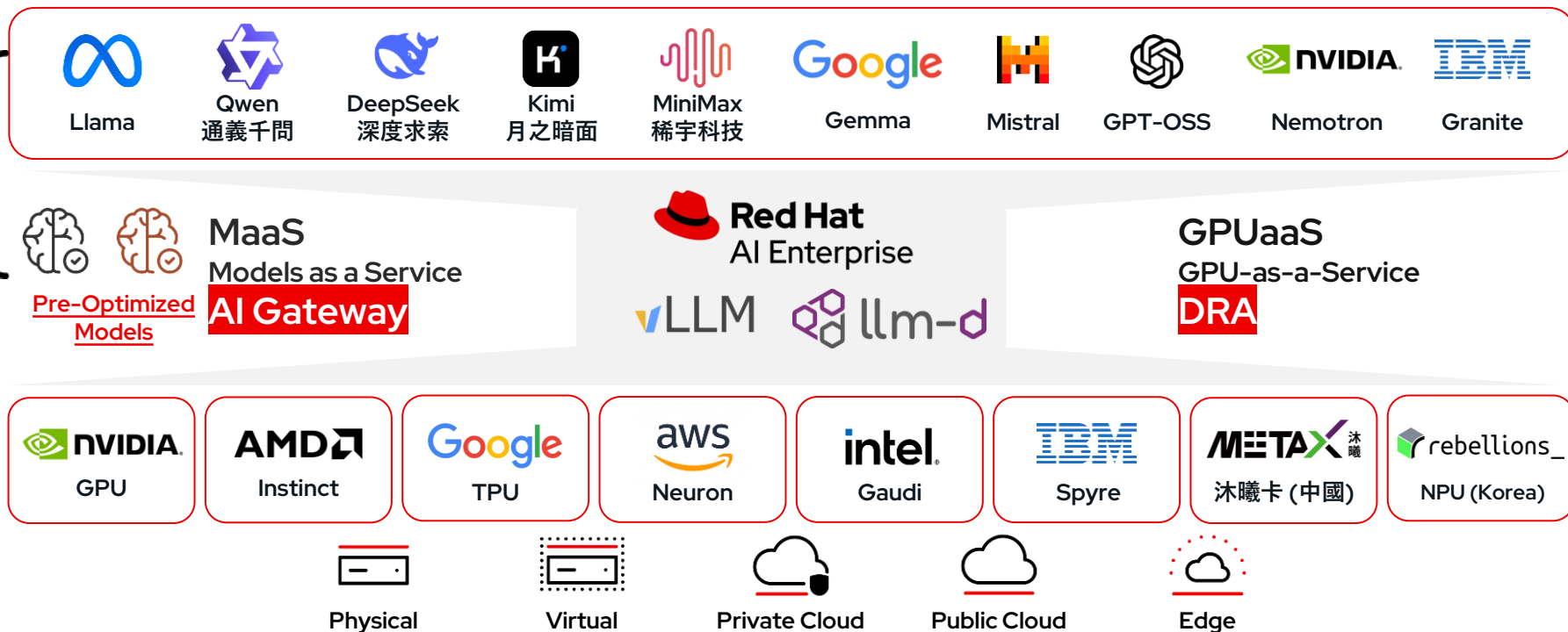
Xeon & CPUs



Spyre



Flexibility from Red Hat AI Enterprise



Single platform to run any model, on any accelerator, on any cloud

Red Hat AI is 100% open source



Thank you!



linkedin.com/company/red-hat



facebook.com/redhatinc



youtube.com/user/RedHatVideos



twitter.com/RedHat

