



# Scalable Inference with llm-d and KServe

Bartosz Majsak

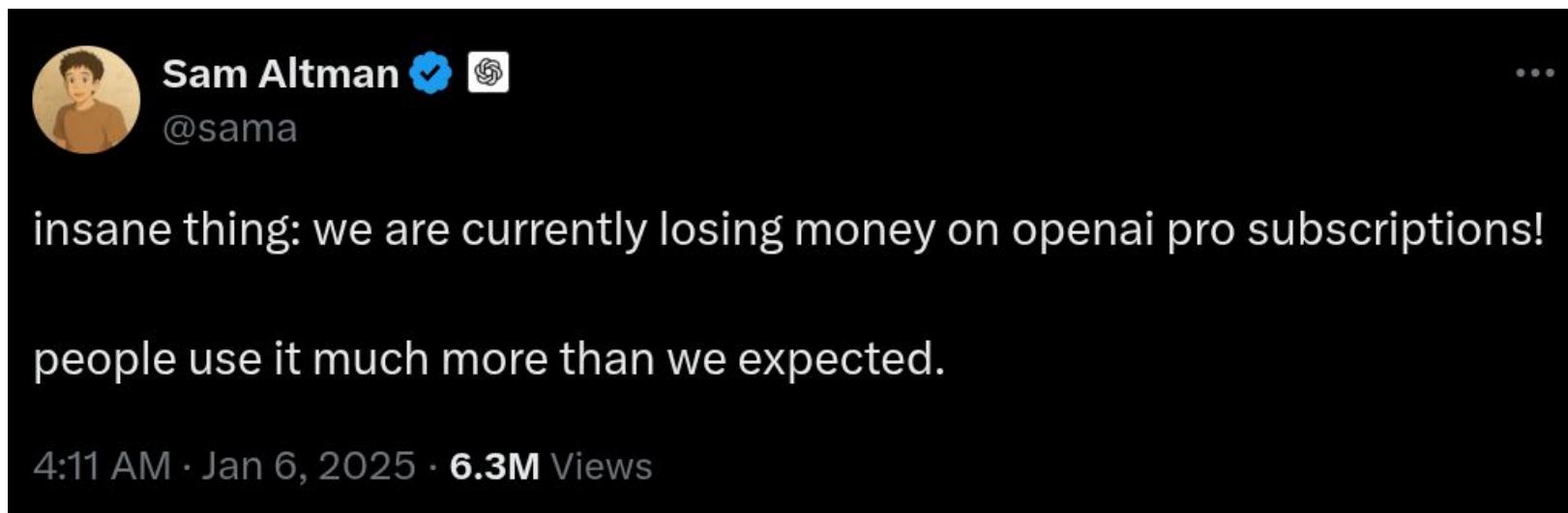


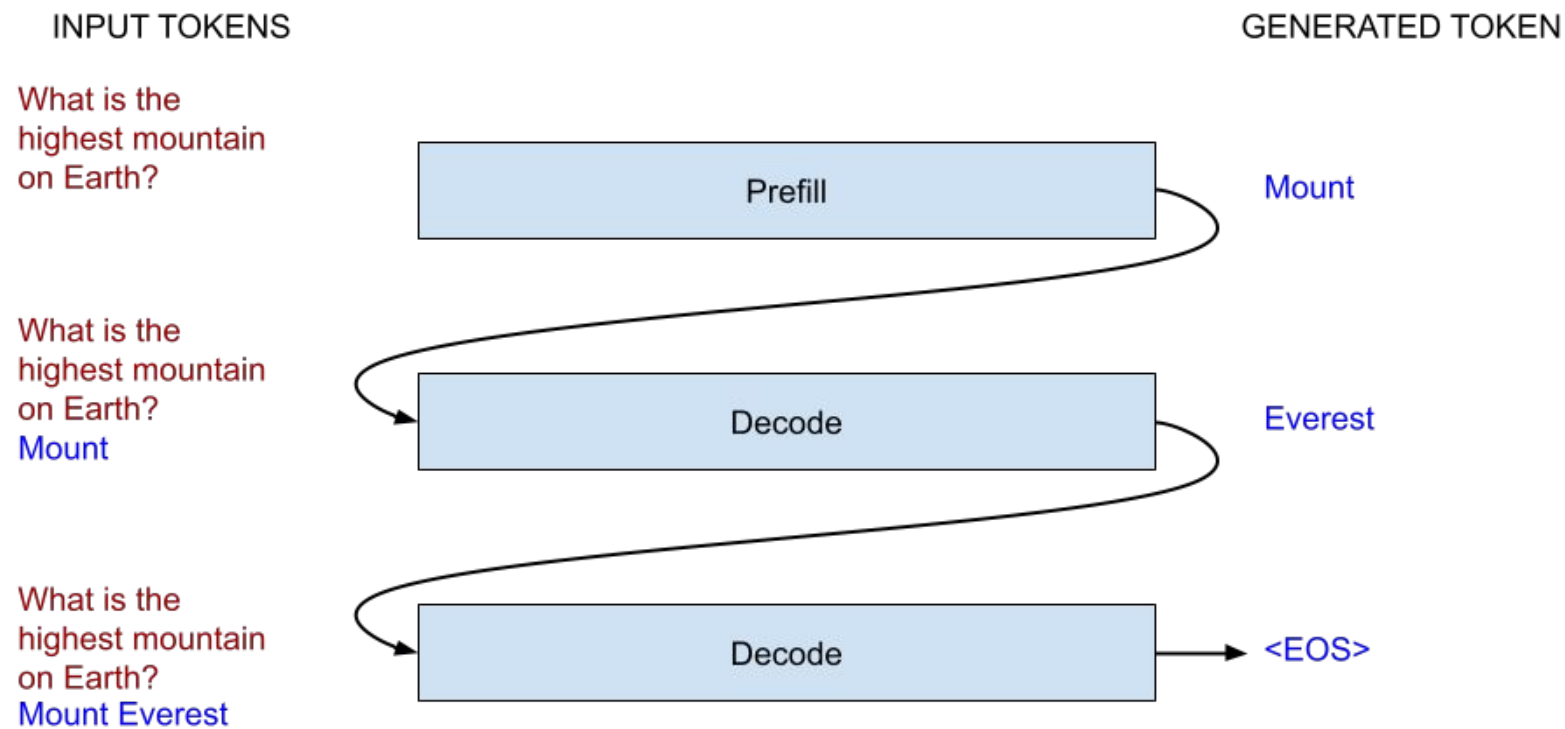
---

# Scaling LLMs is hard

---

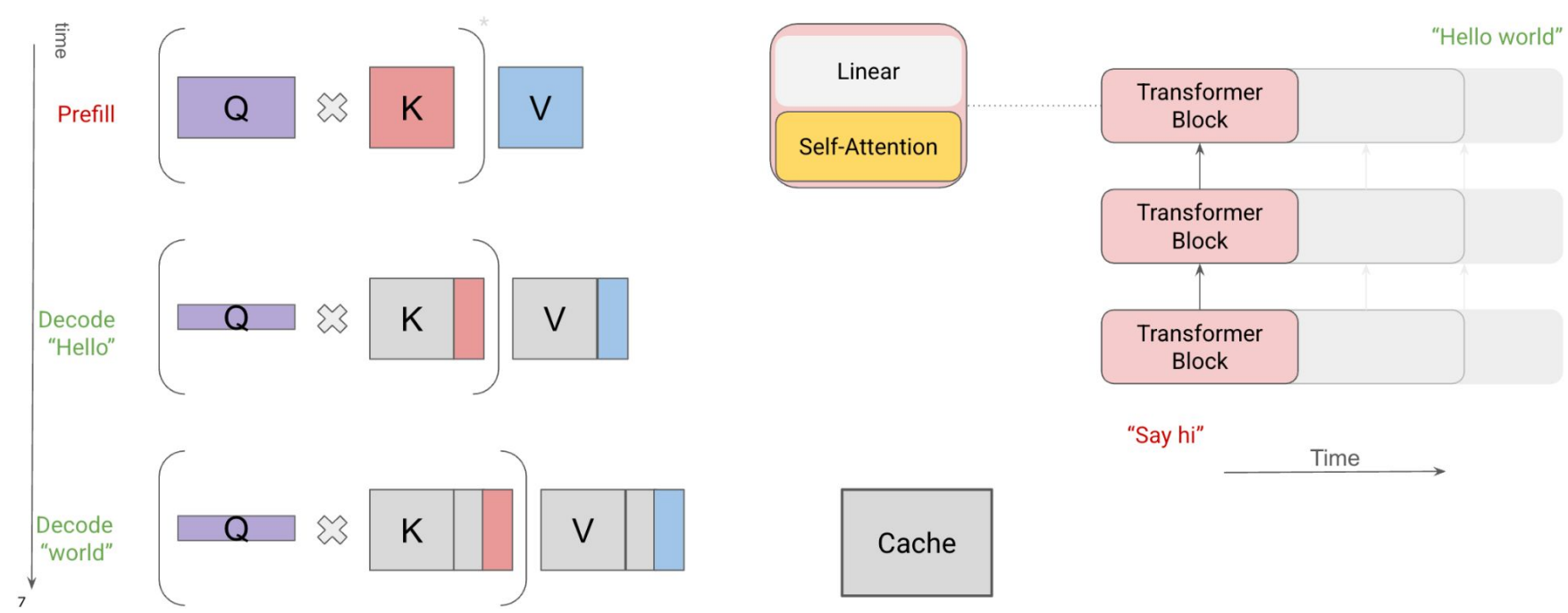
# Scaling LLMs is hard





# KV Cache

Q = What this word "wants to know"  
K = What this word contains  
V = The actual content/information of this word

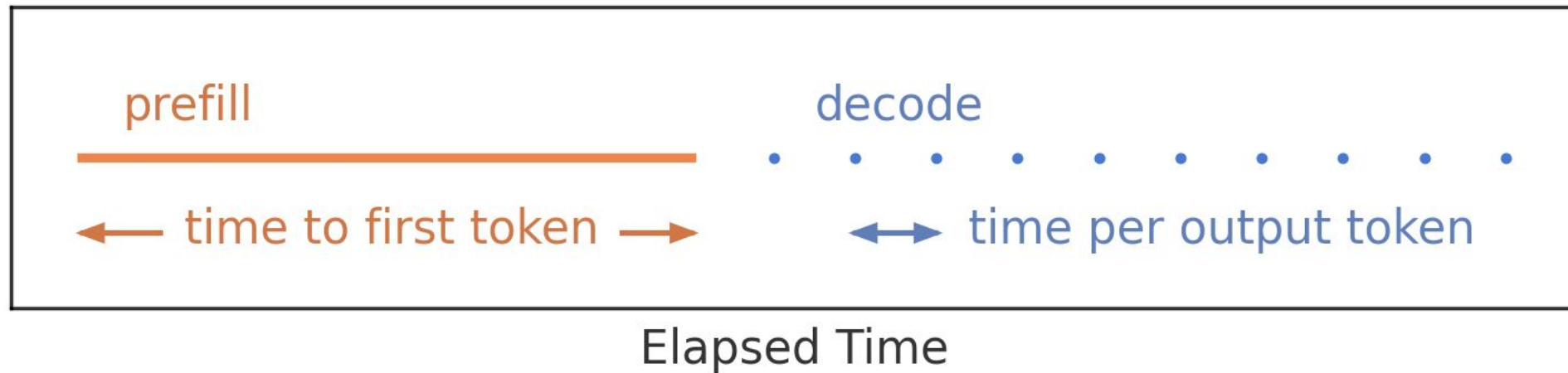


## Prefill

- Processes input tokens and computes
- Highly parallelized, efficient GPU utilization

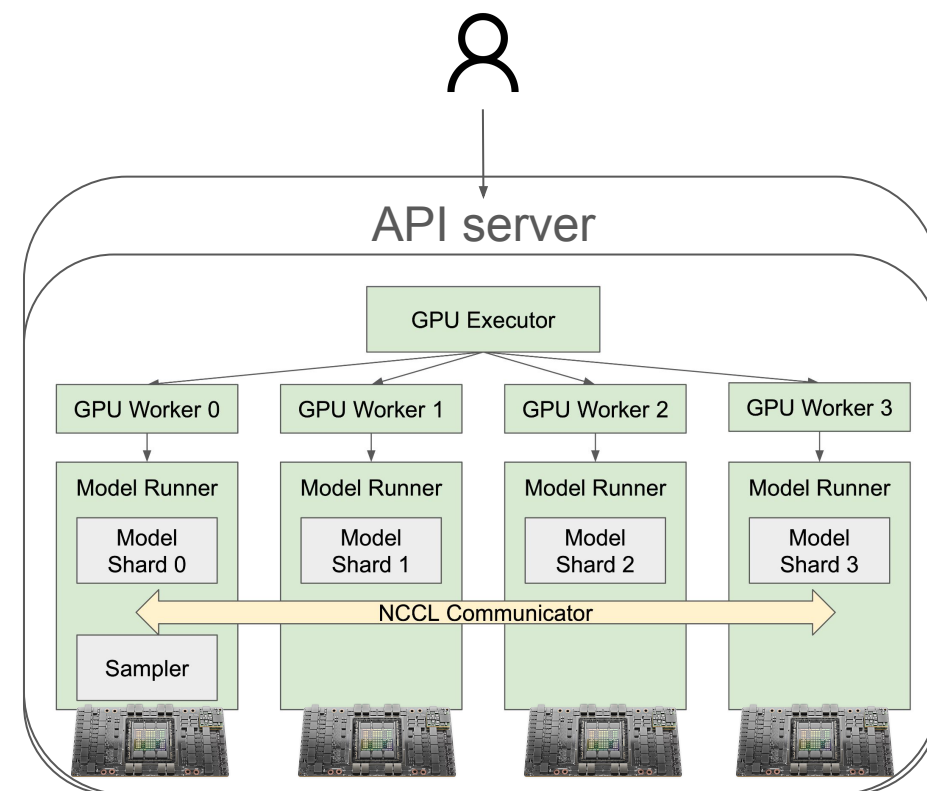
## Decode

- Generates output tokens one at a time
- Sequential processing, based on previous tokens
  - potentially slower
  - memory-bound



# : a LLM Inference & Serving Engine

- ▶ Library or REST API (OpenAI API specification)
- ▶ Allowing users to chat with one model with:
  - **Text, Images, Videos or Audio** as input
  - **Text-only** as output (no diffusion models)
- ▶ **100+** model architectures supported
- ▶ ~400 PR/months
- ▶ Supports “all” GPUs/accelerators
  - NVIDIA, AMD, Intel, IBM Spyre, Google TPU, AWS Trainium/Inferentia, etc...

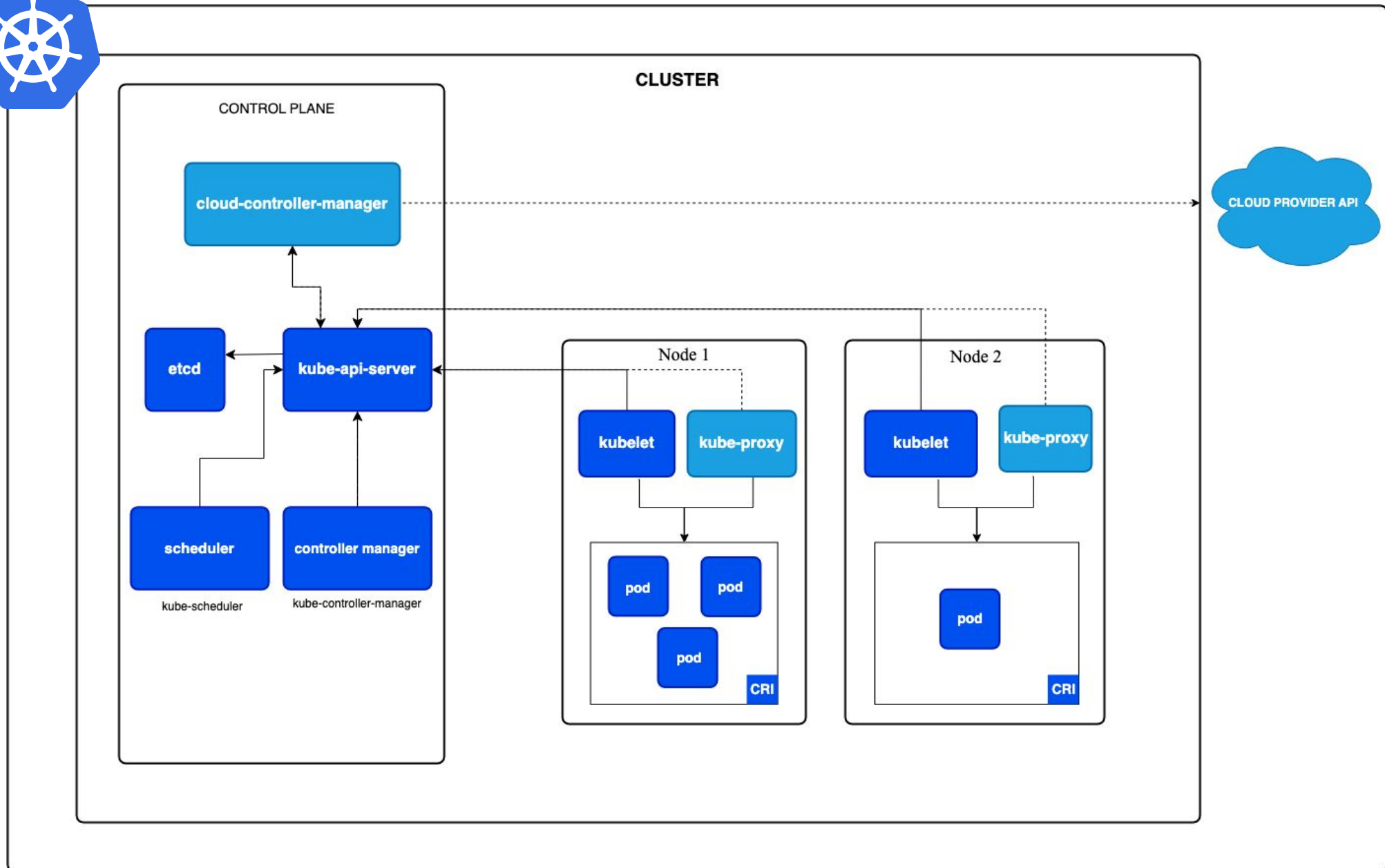
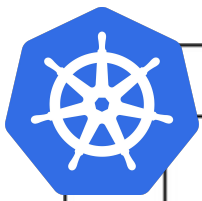


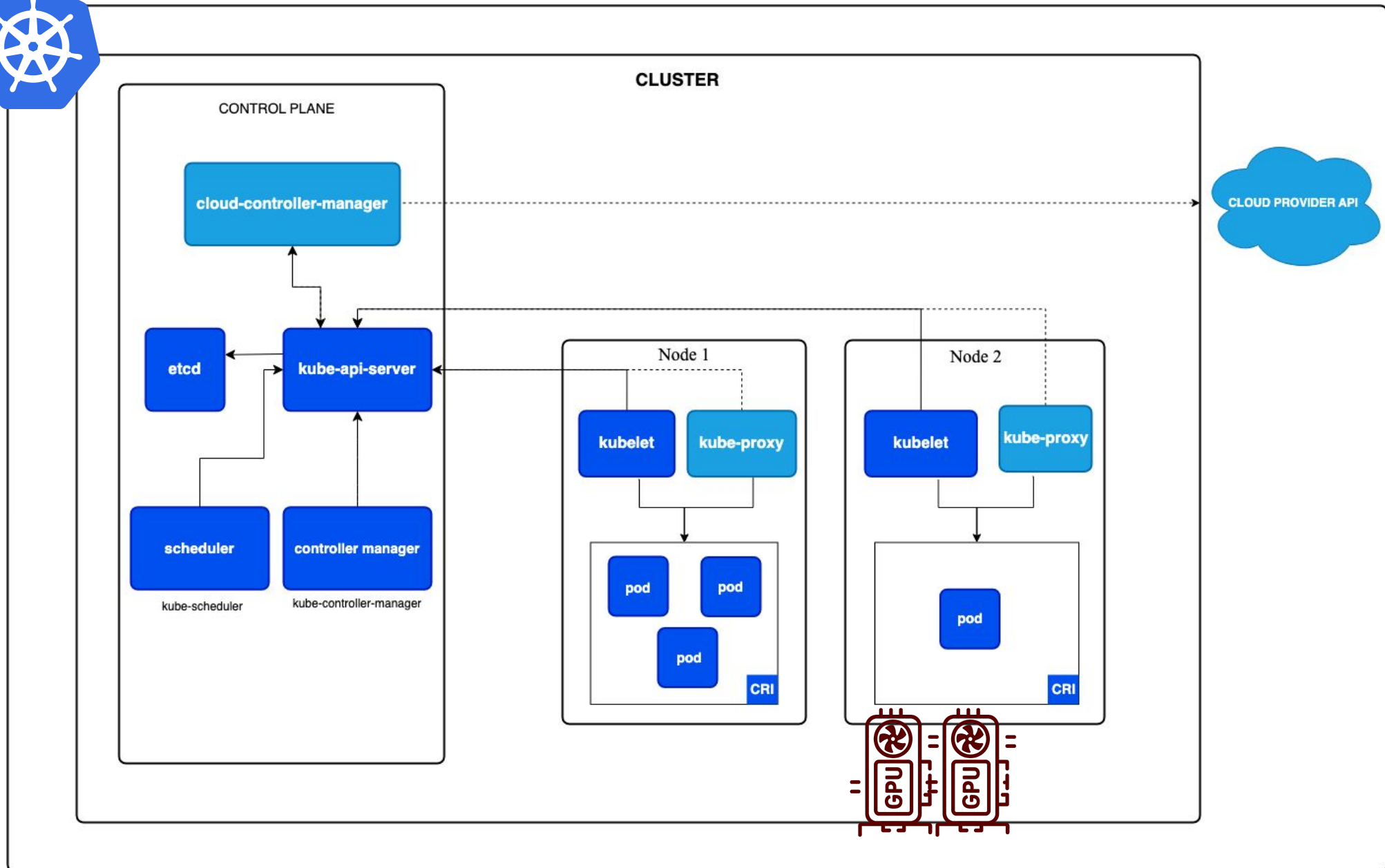
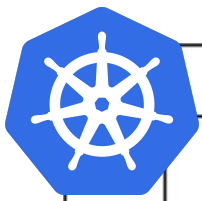
---

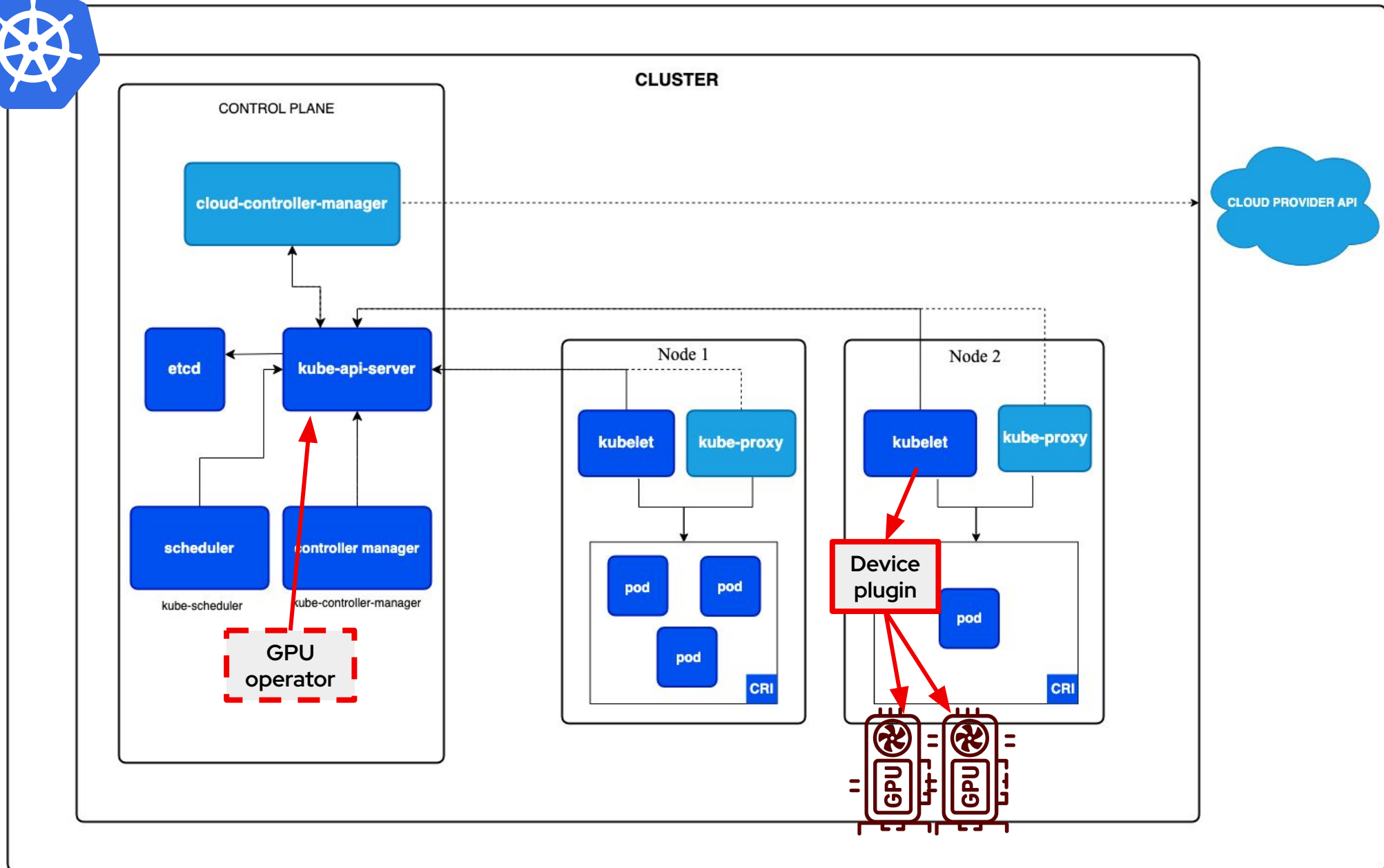
# Kubernetes GPU management

**CONTAINERIZE**









I need GPU!

```
spec:
  containers:
  - name: app
    image: ...
    resources:
      requests:
        memory: "64Mi"
        cpu: "250m"
        nvidia.com/gpu: 2
      limits:
        memory: "128Mi"
        cpu: "500m"
```



# But... a node with GPU is expensive

Don't waste it!

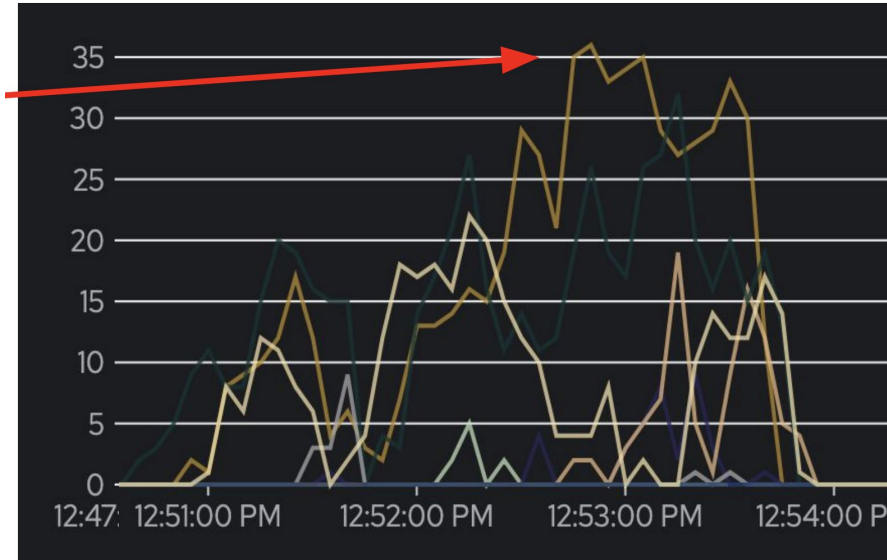
```
apiVersion: machine.openshift.io/v1beta1
kind: MachineSet
metadata:
  ...
spec:
  replicas: 1
  selector:
    ...
  template:
    ...
    spec:
      ...
      taints:
        - key: restrictedaccess
          value: "yes"
          effect: NoSchedule
```

```
apiVersion: nvidia.com/v1
kind: ClusterPolicy
metadata:
  ...
  name: gpu-cluster-policy
spec:
  ...
  daemonsets:
    ...
    tolerations:
      - effect: NoSchedule
        key: restrictedaccess
        operator: Exists
  sandboxWorkloads: ...
  gds: ...
  vgpuManager: ...
  vfioManager: ...
  toolkit: ...
  ...
```

---

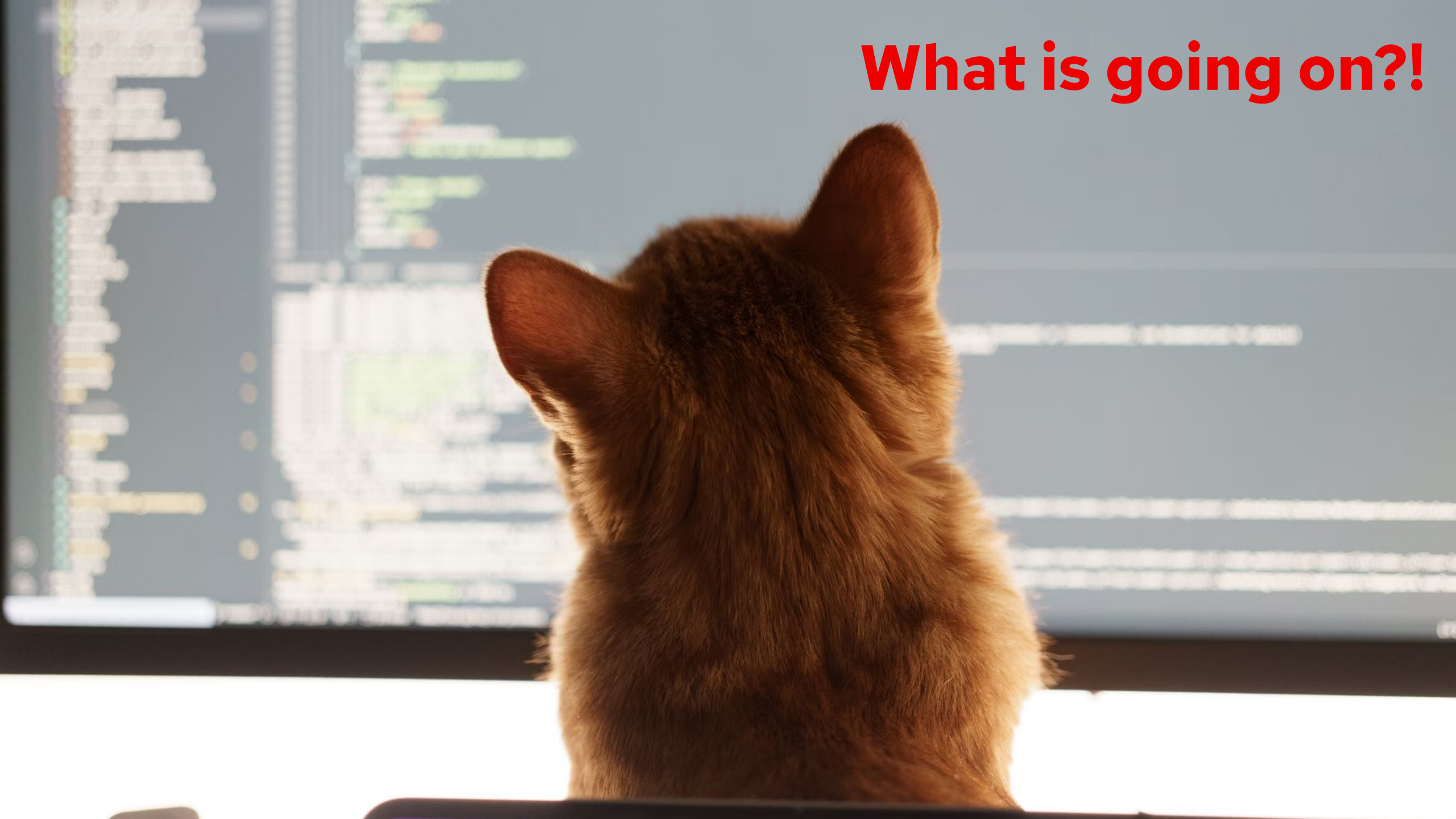
# Disaggregated serving

# When a single instance is not enough...



vLLM metric: num\_requests\_waiting  
(raw-deployment)

**What is going on?!**



# Prefill and decode

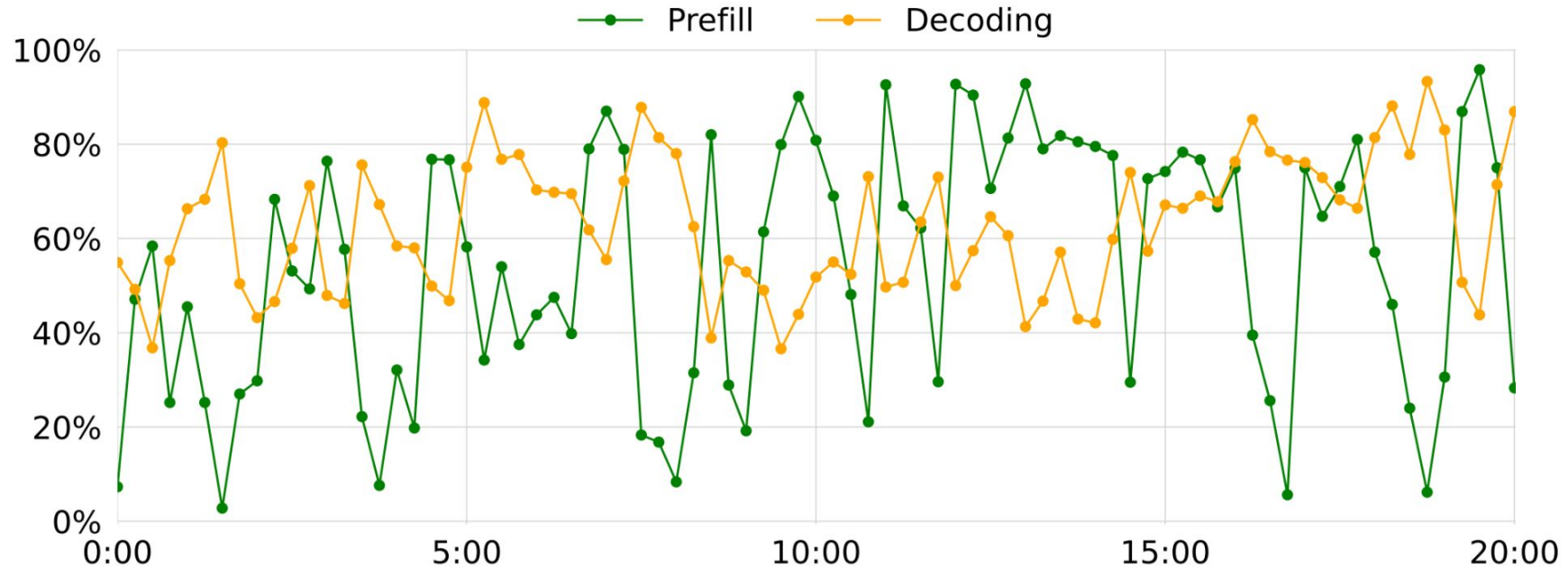


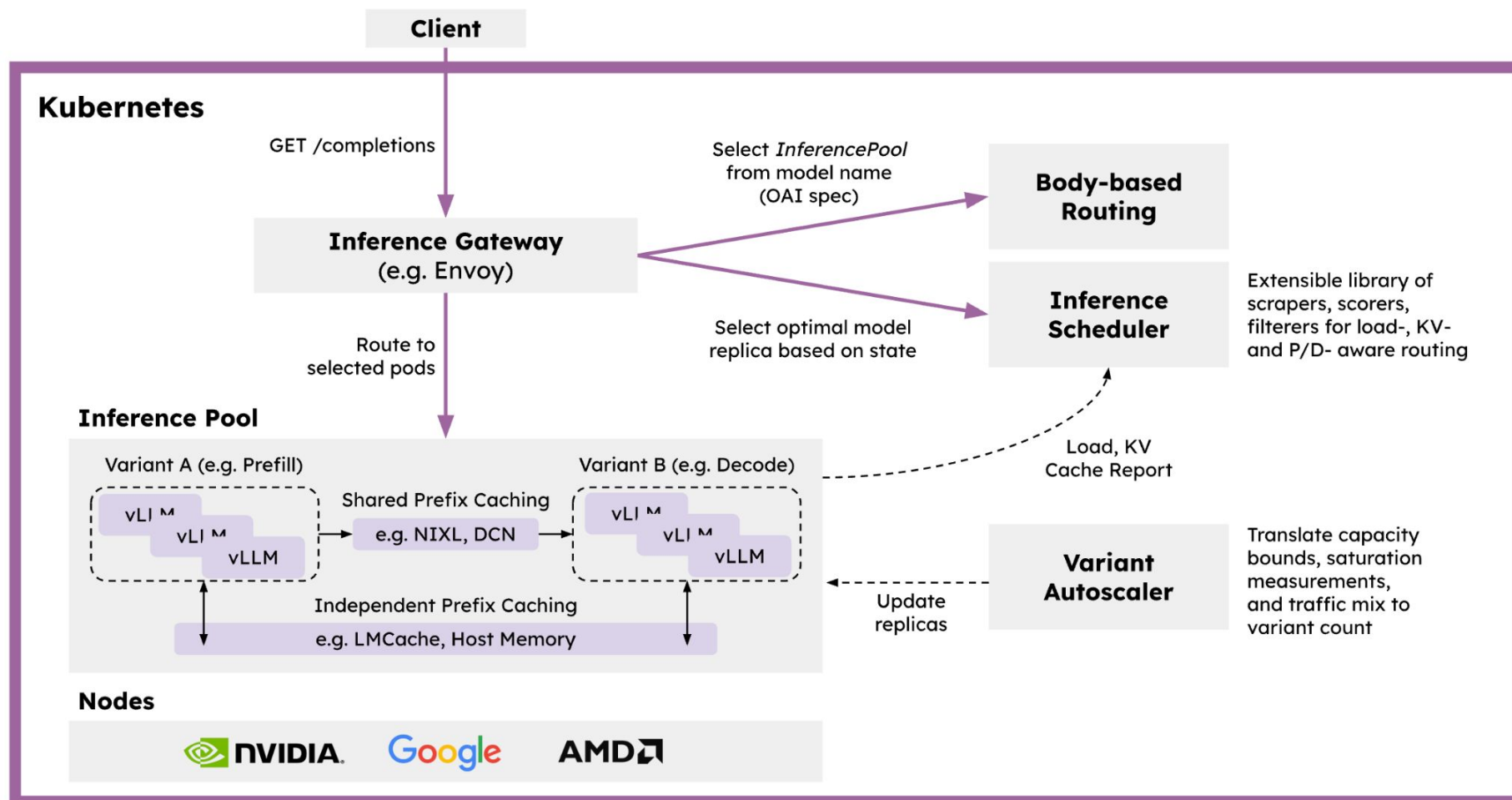
Figure 9: The load of prefill and decoding instances over 20 minutes, before using the prediction-based early rejection.

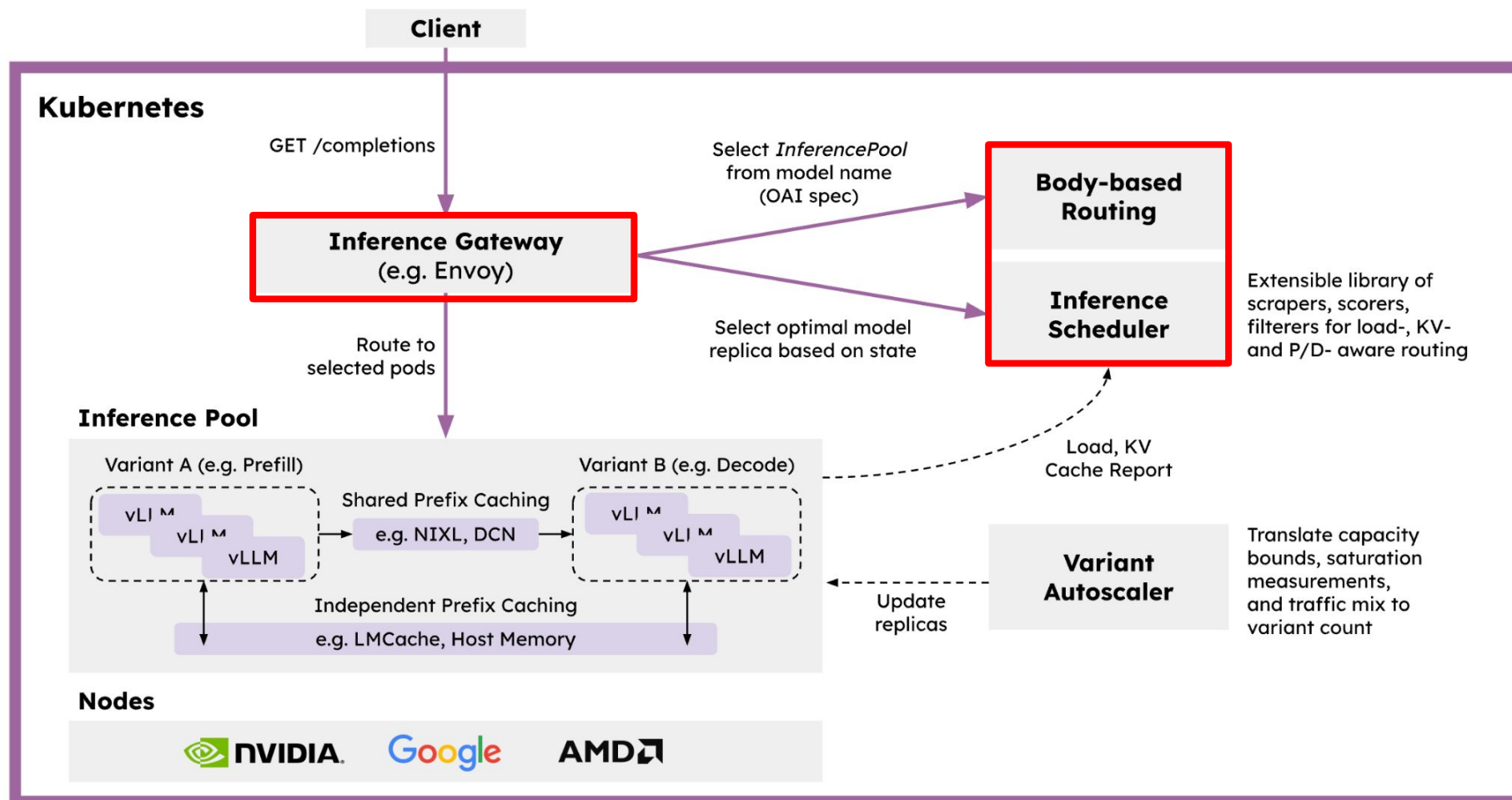
# What is going on? – Prefill and decode

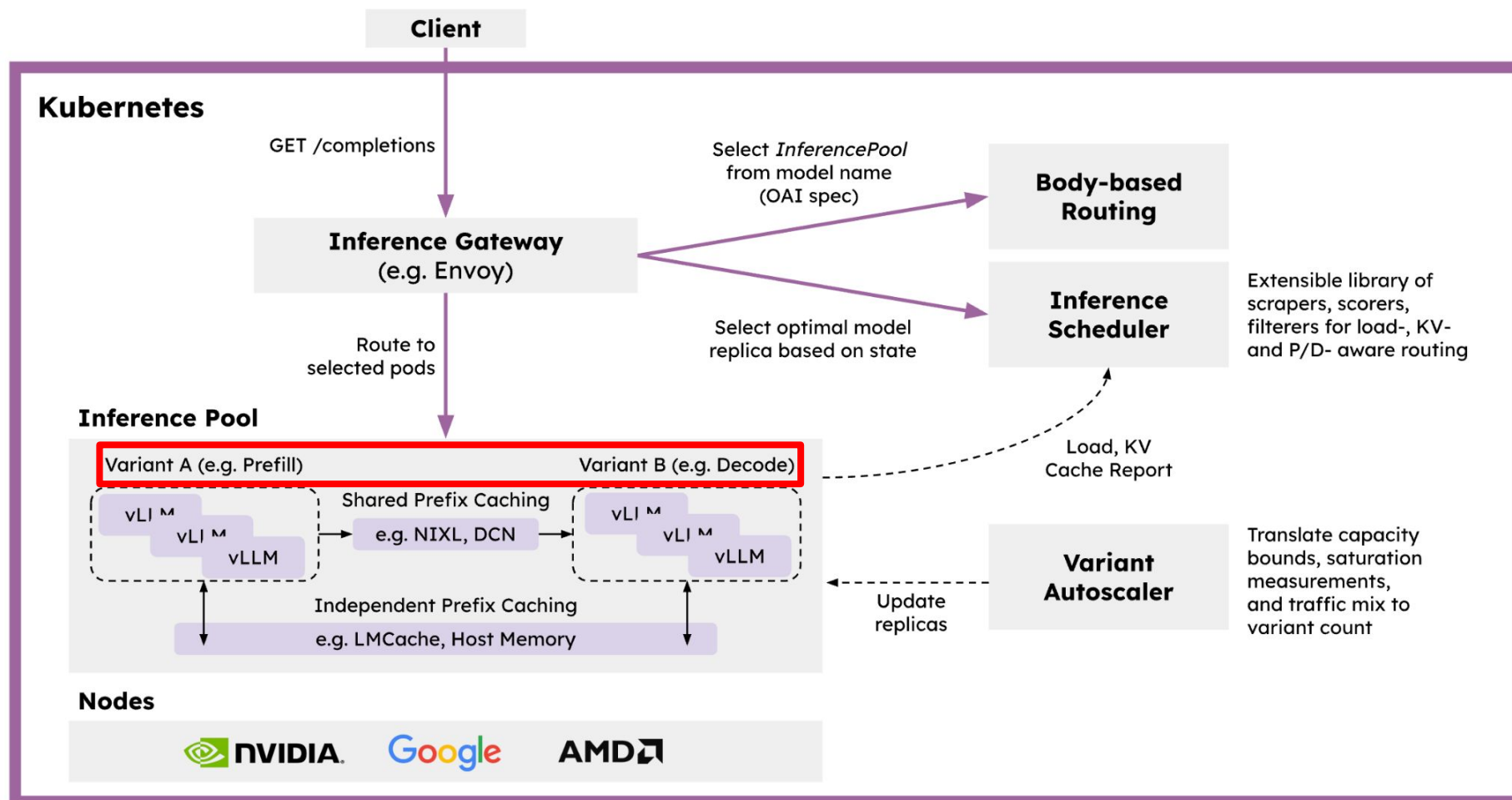
*All ~~animals~~ requests are **NOT** equal*

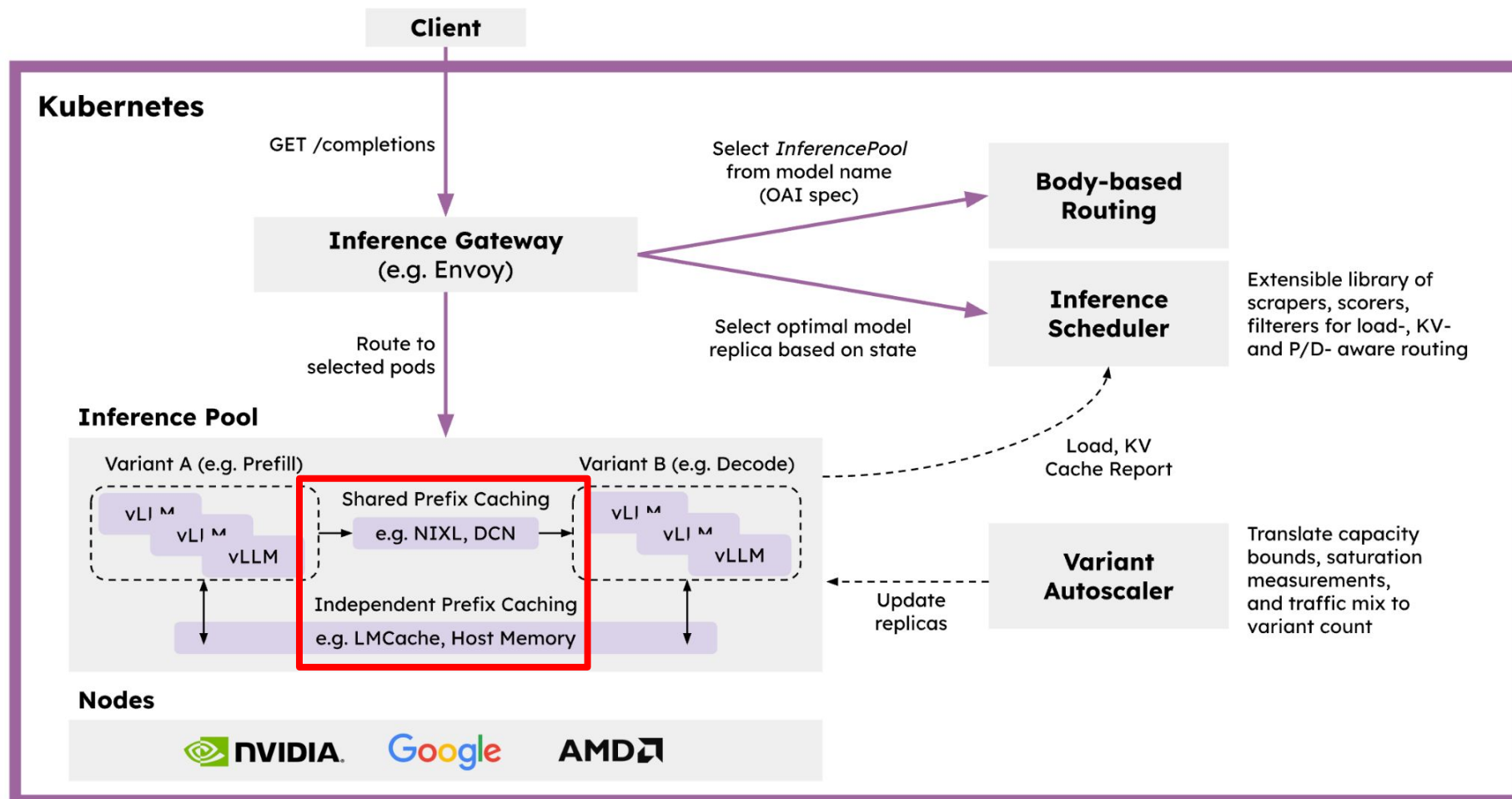
- ▶ **Prefill** and **Decode** phases have **different scalability challenges** and latency implication.
- ▶ **Prefill phase** processes the prompt produces the first token (**Time To First Token**)
- ▶ **Decode phase** produces all the other tokens (**Inter Token Latency**)

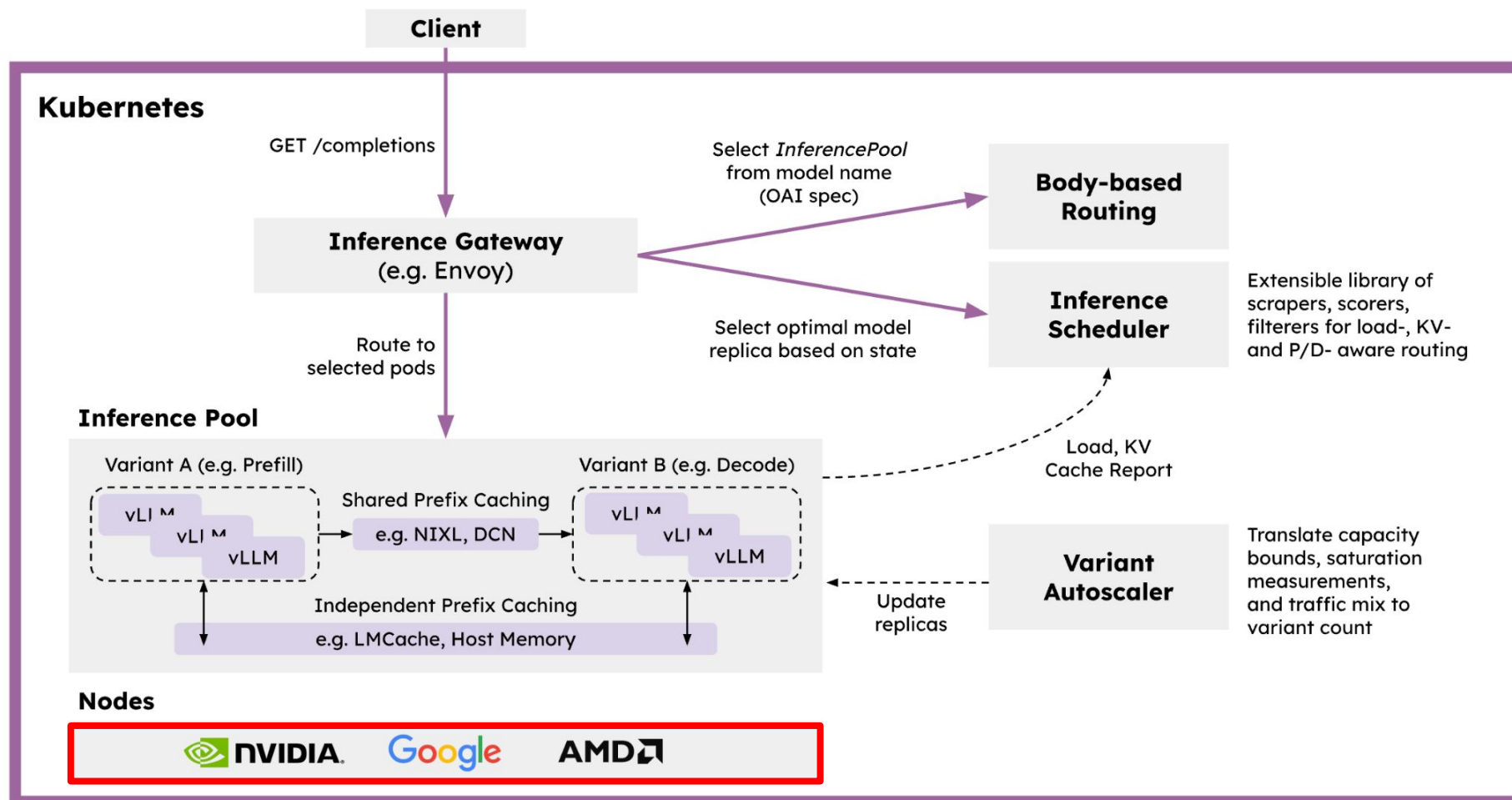


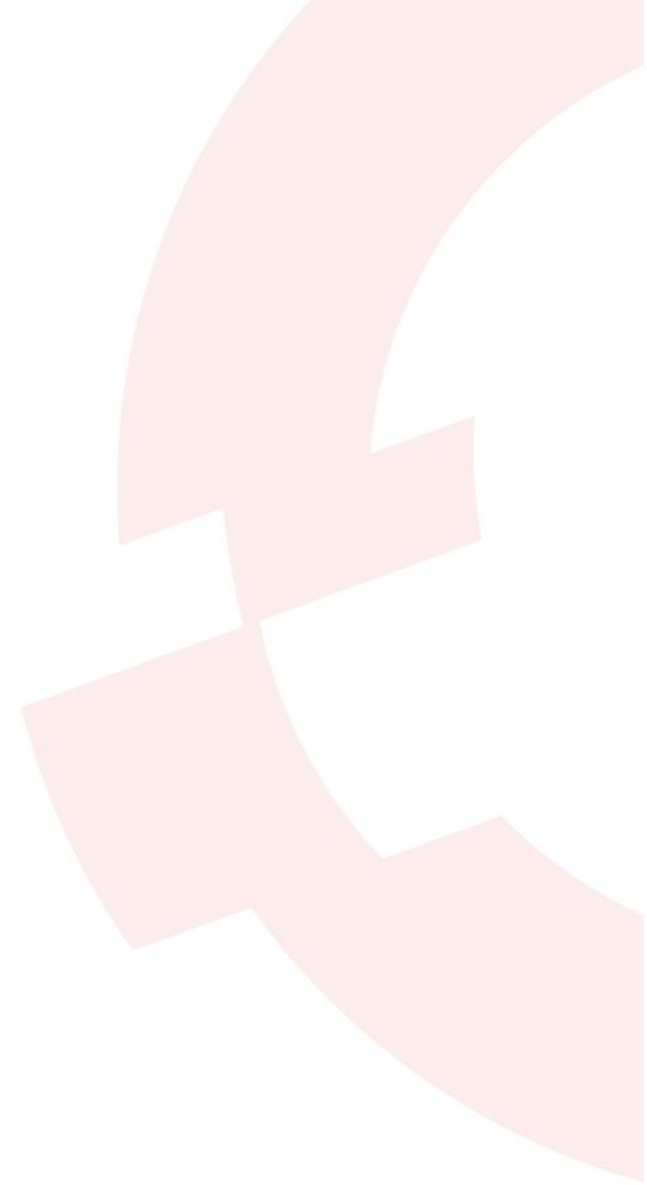


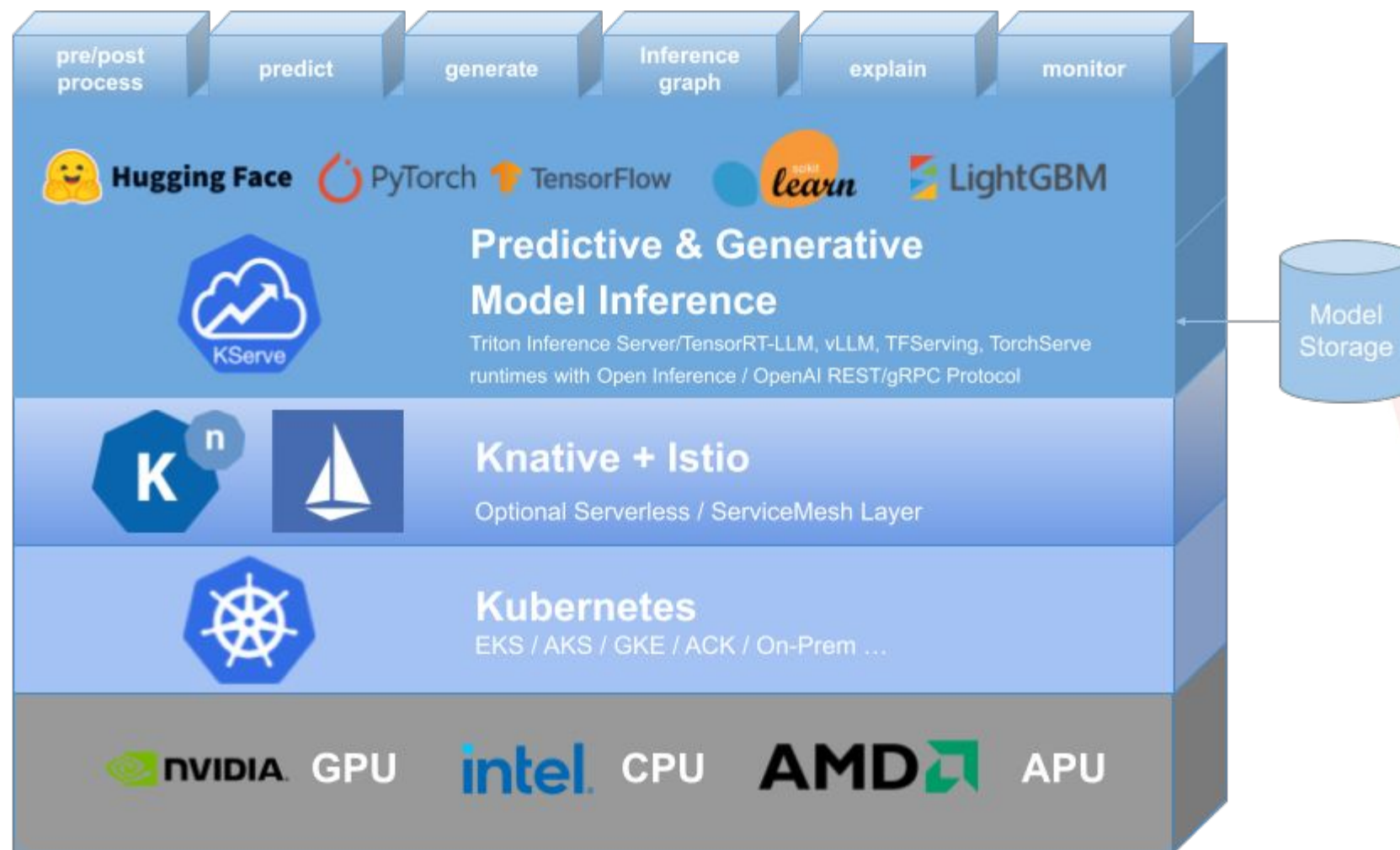












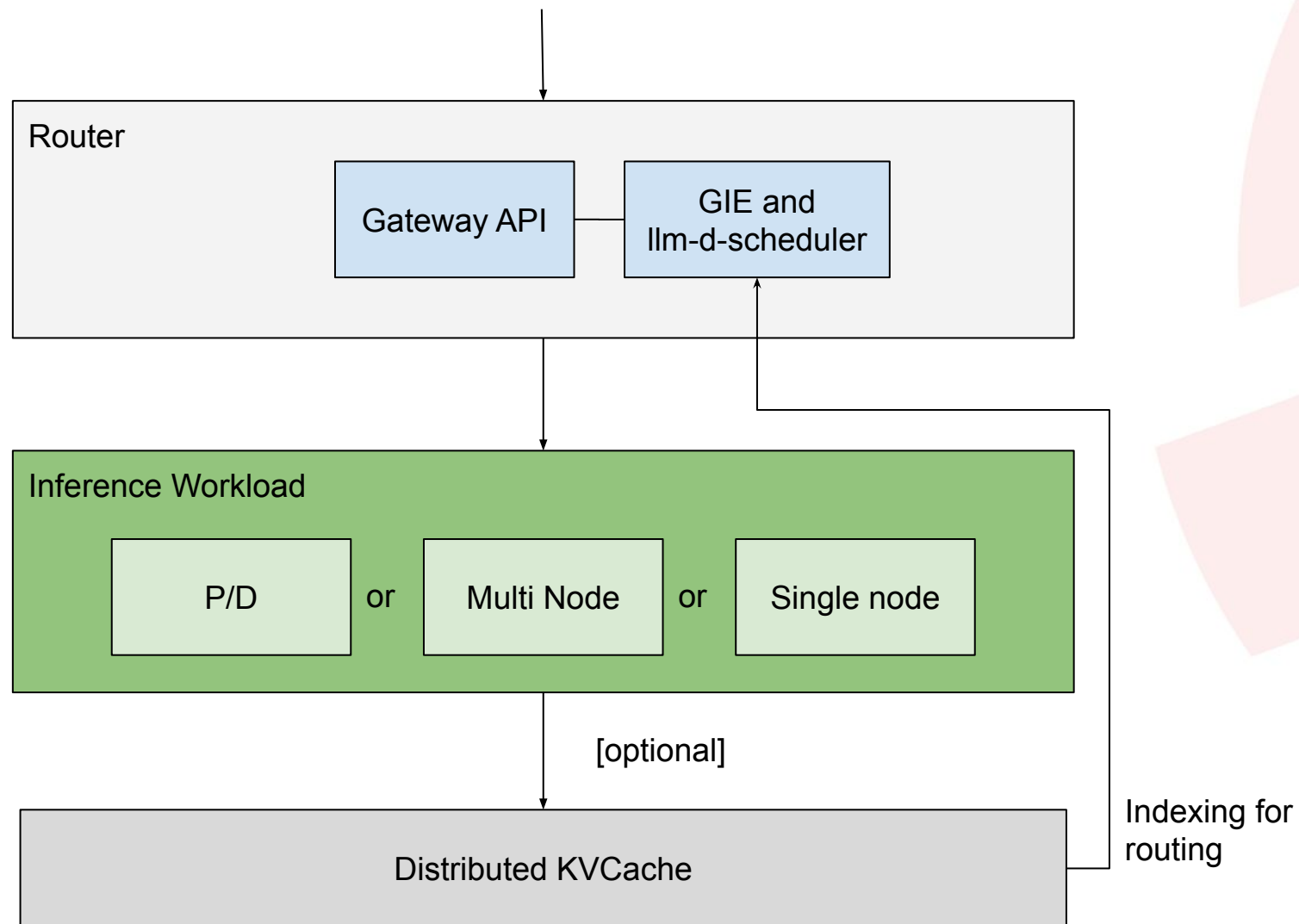
```
kind: LLMInferenceService
metadata:
  name: llama4-scout-17b-16e-instruct
spec:
  replicas: 2
  model:
    uri: hf://meta-llama/Llama-4-Scout-17B-16E-Instruct
  router: # automatically create Gateway and HTTPRoute
    route: {}
    gateway: {}
    scheduler: {}
```

```
kind: LLMInferenceServiceConfig
metadata:
  name: nvidia-a100
spec:
  template:
    nodeSelector:
      nvidia.com/gpu.product: NVIDIA-A100-PCIE-40G
    containers:
      - name: main
        image: vllm-image...
        resources:
          requests:
            nvidia.com/gpu: "1"
```

```
kind: LLMInferenceService
metadata:
  name: llama4-scout-17b-16e-instruct
spec:
  replicas: 2
  model:
    uri: hf://meta-llama/Llama-4-Scout-17B-16E-Instruct
  router:
    route: {}
    gateway: {}
    scheduler: {}
  baseRefs:
    - nvidia-a100
```

```
spec:
  router: # ...
```

```
template: # ...
worker:   # ...
prefill:  # ...
```





Demo!

## Helpful links

- <https://kserve.github.io/website/>
- <https://llm-d.ai/>
- <https://docs.vllm.ai/en/stable/>
- <https://ai-on-openshift.io/>
- <https://www.redhat.com/de/products/ai/openshift-ai>
- <https://gateway-api.sigs.k8s.io/>
- <https://gateway-api-inference-extension.sigs.k8s.io/>

# Thank you

Red Hat is the world's leading provider of enterprise open source software solutions. Award-winning support, training, and consulting services make Red Hat a trusted adviser to the Fortune 500.



[linkedin.com/company/red-hat](https://www.linkedin.com/company/red-hat)



[facebook.com/redhatinc](https://www.facebook.com/redhatinc)



[youtube.com/user/RedHatVideos](https://www.youtube.com/user/RedHatVideos)



[twitter.com/RedHat](https://twitter.com/RedHat)