

Technical Case Studies

In-depth technical case studies exploring real-world architecture decisions, implementation challenges, and engineering solutions from production systems.

Contents

01	Netflix's In-House LLM Serving: Technical Case Study	3
-----------	--	---

Netflix's In-House LLM Serving: Technical Case Study

Executive Summary

Netflix, a global leader in streaming entertainment, has strategically invested in developing and deploying its own Large Language Models (LLMs) to power a new generation of personalized experiences and content creation workflows. This technical case study delves into their in-house LLM serving architecture, a complex system designed to meet stringent demands for low-latency inference, high throughput, cost efficiency, and data privacy. By building a dedicated MLOps platform for LLMs, Netflix addresses the unique challenges of managing multi-billion parameter models, from dynamic resource allocation on GPU clusters to sophisticated model versioning and real-time observability. This analysis highlights key architectural decisions, the trade-offs encountered under production loads, and the practical solutions implemented, offering valuable insights for platform engineers and AI developers embarking on similar endeavors.

The Netflix Context: Why In-House LLMs?

Netflix's business thrives on hyper-personalization and efficient content production. While external LLM providers offer convenience, the unique scale, data sensitivity, and specific customization needs of Netflix necessitate an in-house solution.

 **Key Idea:** Control over data, cost, and customization drives strategic in-house LLM development.

- **Data Sovereignty and Privacy:** User viewing habits, content metadata, and internal communication are highly sensitive. Relying on third-party LLM APIs introduces data governance and privacy risks that conflict with Netflix's strict policies.
- **Hyper-Personalization at Scale:** Generic LLMs often fall short in understanding the nuanced context of Netflix's content catalog and diverse global audience. In-house models can be fine-tuned on proprietary data, leading to more relevant recommendations, search results, and generative content.

- **Cost Efficiency for High Volume:** At Netflix's scale, external API calls for millions of daily inference requests can quickly become prohibitively expensive. Operating custom models on optimized infrastructure offers significant long-term cost savings.
- **Innovation and Differentiation:** Building and serving LLMs internally fosters deeper research, enables rapid experimentation with novel architectures, and unlocks unique capabilities that differentiate Netflix's product experience. This includes advanced content tagging, script analysis, localization, and interactive storytelling tools.

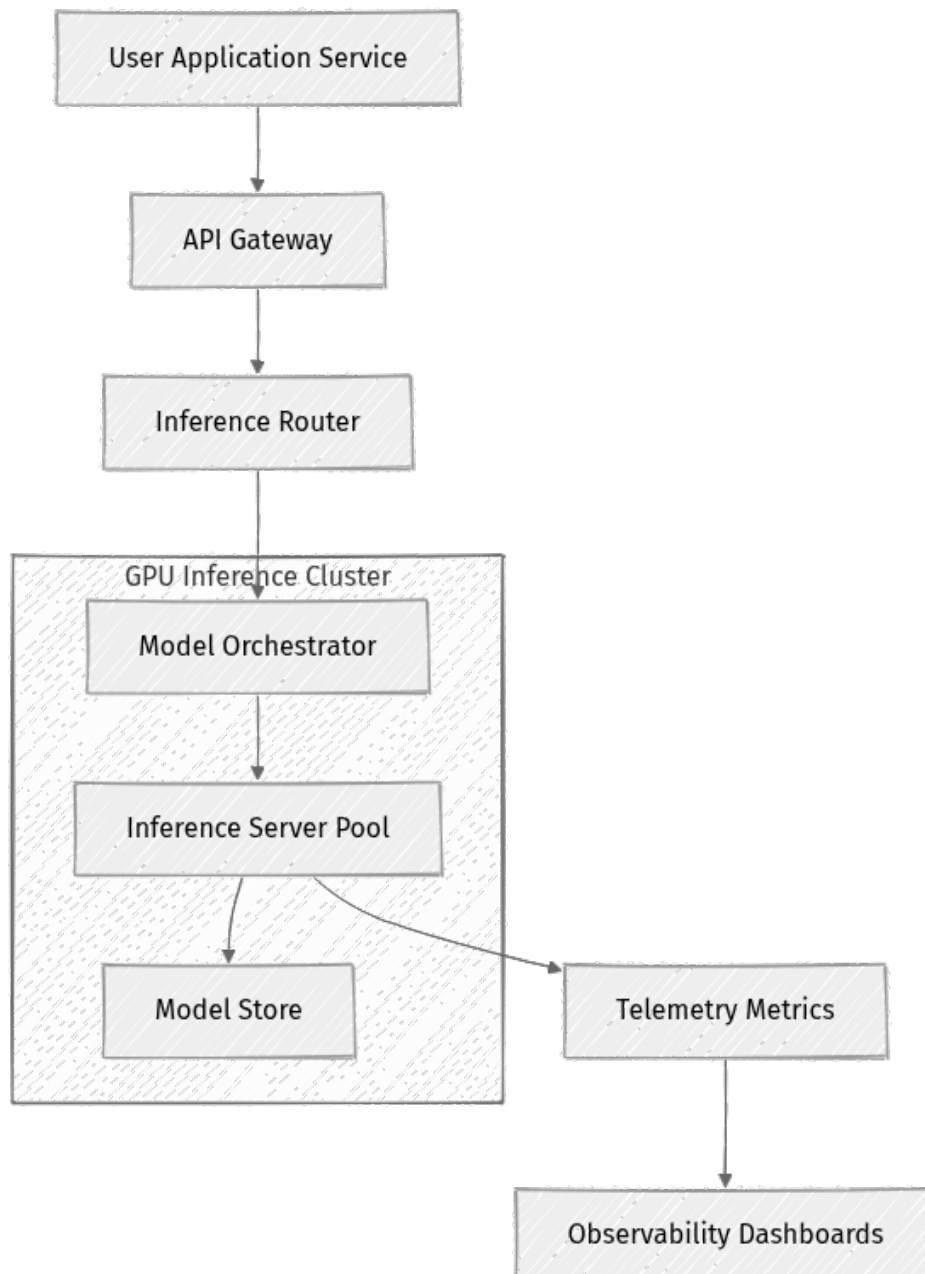
Core Requirements for LLM Inference

Deploying LLMs into production at Netflix's scale imposed a demanding set of technical requirements that shaped the entire serving architecture.

- **Low Latency Inference:** For user-facing applications like real-time recommendations, interactive search, or personalized content summaries, inference must complete within tens to hundreds of milliseconds.
- **High Throughput:** The system must handle millions of inference requests per second globally, supporting diverse applications simultaneously.
- **Cost Optimization:** GPU resources are expensive. The architecture must maximize GPU utilization, dynamically scale resources, and leverage cost-effective hardware strategies.
- **Model Heterogeneity:** Support for a wide range of LLM architectures, sizes (from few-billion to hundreds-of-billions of parameters), and fine-tuned variants.
- **Scalability and Elasticity:** The system needs to scale up and down rapidly to meet fluctuating demand, from peak viewing hours to off-peak periods.
- **Reliability and Resilience:** High availability with robust failover mechanisms and graceful degradation under extreme load.
- **Observability:** Comprehensive monitoring of model performance, infrastructure health, resource utilization, and inference costs.
- **Security:** Isolation of data and models, secure API endpoints, and strict access controls.

Architectural Blueprint: A Multi-Layered Approach

Netflix's in-house LLM serving architecture is a sophisticated, multi-layered system designed for resilience, performance, and scalability. It integrates seamlessly with their existing cloud-native infrastructure, primarily built on AWS and their internal container orchestration platform, Titus (a custom extension of Kubernetes).



Architecture Components:

- **User Application/Service:** Any internal or external-facing Netflix service that requires LLM inference (e.g., recommendation engines, content creation tools, customer support bots).
- **API Gateway:** Acts as the entry point for all LLM inference requests, handling authentication, authorization, rate limiting, and request routing. This leverages Netflix's existing Edge infrastructure.
- **Inference Router:** A specialized service responsible for directing incoming requests to the appropriate LLM endpoint. It considers factors like model version, regional deployment, and current load. This layer might also handle request queuing and dynamic batching logic.
- **Model Orchestrator:** The brain of the serving system. It manages the lifecycle of LLM instances on the GPU cluster. This includes:
 - **Dynamic Scaling:** Spinning up/down inference servers based on demand.
 - **Model Placement:** Deciding which models run on which GPU instances.
 - **Health Checks:** Monitoring the status of inference servers.
 - **Version Management:** Facilitating blue/green deployments and rollbacks for models.
- **GPU Inference Cluster:** The core compute layer, consisting of a fleet of GPU-accelerated instances (e.g., AWS EC2 P-series or G-series). Managed by Titus for container orchestration.
- **Inference Server Pool:** Individual containers running specialized LLM inference engines (e.g., vLLM, NVIDIA Triton Inference Server with TensorRT-LLM backend). Each server can host one or multiple models, often leveraging techniques like multi-model serving or continuous batching.
- **Model Store:** A highly available, low-latency object storage system (e.g., S3-compatible storage) where all LLM model artifacts (weights, configurations, tokenizers) are stored and versioned. Inference servers pull models from here.
- **Telemetry/Metrics:** A comprehensive system for collecting real-time metrics (latency, throughput, error rates, GPU utilization, memory consumption, cost per inference) and logs from all components.

- **Observability Dashboards:** Visualization and alerting tools (e.g., Netflix Atlas, custom dashboards) to provide operational insights and enable rapid incident response.

Key Infrastructure Pillars: Model Storage, Serving, and Orchestration

The success of Netflix's LLM serving hinges on robust implementations across three core infrastructure areas:

Model Storage and Versioning

- **Centralized Object Storage:** All LLM artifacts (model weights, tokenizers, configuration files, quantization tables) are stored in a highly available, S3-compatible object storage system. This ensures durability, global accessibility, and cost-effectiveness.
- **Immutable Versioning:** Every model update results in a new, immutable version stored in the system. This allows for precise rollbacks, A/B testing, and ensures reproducibility of inference results.
- **Metadata Management:** A metadata catalog tracks details for each model version, including architecture, training data lineage, performance benchmarks, and deployment status. This is crucial for governance and operational visibility.
- **Efficient Distribution:** Models are often very large (hundreds of GBs). To minimize cold start times, models are pre-staged in regional caches or streamed efficiently using techniques like lazy loading of weights.

High-Performance Inference Serving

- **Specialized Serving Runtimes:** Netflix leverages and customizes open-source inference servers optimized for LLMs, such as vLLM for continuous batching and NVIDIA Triton Inference Server with its TensorRT-LLM backend for optimized execution graphs.
- **Dynamic Batching:** To maximize GPU utilization, incoming requests are dynamically grouped into batches. This amortizes the overhead of GPU kernel launches and significantly improves throughput, especially under variable load.

- **Quantization and Sparsity:** Models are often quantized (e.g., to FP8 or INT8) or pruned to reduce their memory footprint and accelerate inference without significant loss in quality. This is a critical step for cost-efficient deployment.
- **Custom Kernels and Optimizations:** For specific model architectures or high-volume operations, Netflix engineers develop custom CUDA kernels or leverage highly optimized libraries (e.g., FlashAttention) to push performance boundaries.

Orchestration and Resource Management with Titus

Netflix's container orchestration platform, Titus (based on Kubernetes principles), is central to managing the dynamic nature of LLM serving.

- **GPU-Aware Scheduling:** Titus is extended with GPU-specific scheduling capabilities, allowing it to efficiently allocate GPU resources (e.g., specific GPU types, memory, compute units) to inference server pods.
- **Auto-Scaling Policies:** Horizontal Pod Autoscalers (HPAs) are configured based on metrics like request queue length, GPU utilization, and latency targets. This enables elastic scaling of inference server pools.
- **Blue/Green Deployments:** New model versions or inference server configurations are deployed using blue/green strategies to ensure zero-downtime updates and easy rollbacks. Traffic is gradually shifted to the new 'green' environment after validation.
- **Resource Isolation:** Each inference server runs in an isolated container, preventing resource contention and ensuring stability across different models and workloads.

```
# Example (simplified) Titus Job definition for an LLM Inference Server
# This is illustrative, actual configs are more complex and dynamic.
name: llm-inference-service-v1_2
jobType: service
owner: ml-platform-team
version: 1.2.0

containers:
- name: llm-server
  image: netflix/llm-inference-base:vllm-0.5.1-cuda12.2
  cpu: 8
  memory: 64Gi
  gpu: 1 # Request 1 GPU
  ports:
    - name: http
      containerPort: 8000
  env:
    MODEL_NAME: "netflix-llama-70b-v1.2"
    MODEL_PATH: "s3://netflix-llm-models/llama-70b/v1.2/"
```

```
QUANTIZATION: "fp8"
MAX_BATCH_SIZE: "128"
command: ["/usr/local/bin/vllm_server.sh"]
args: ["--model", "$MODEL_PATH", "--quantization", "$QUANTIZATION"]

resourcePool: gpu-inference-pool
scaling:
  minInstances: 2
  maxInstances: 20
  policy:
    type: METRIC_BASED
    metrics:
      - name: RequestQueueDepth
        targetValue: 50
      - name: GpuUtilization
        targetValue: 0.8 # Scale up if GPU utilization exceeds 80%
```

Optimizing for Production Load: Latency and Throughput

Achieving Netflix's performance targets for LLM inference requires relentless optimization at every layer.

- **Continuous Batching (vLLM):** Unlike traditional batching that waits for a full batch, continuous batching processes tokens as soon as they are generated, maximizing GPU utilization by overlapping computation with I/O and handling variable sequence lengths efficiently. This is a game-changer for high-throughput, low-latency scenarios.
- **Flash Attention:** A highly optimized attention mechanism that reduces memory usage and speeds up computation for long sequences, crucial for complex prompts or multi-turn conversations.
- **Speculative Decoding:** For certain generative tasks, a smaller, faster "draft" model generates initial tokens, which a larger "verifier" model then quickly checks. This can significantly speed up token generation.
- **Model Partitioning and Pipeline Parallelism:** For extremely large models that don't fit on a single GPU, the model can be split across multiple GPUs or even multiple nodes. Pipelining techniques allow different GPUs to work on different layers simultaneously.
- **Caching Layers:**
 - **Prompt Caching:** Caching the output of common or identical prompts to avoid re-computation.
 - **KV Cache Management:** Efficiently managing the key-value cache (attention states) on the GPU to reduce memory pressure, especially for long-running conversations.

- **Hardware Selection:** Continuous evaluation and adoption of the latest GPU architectures (e.g., NVIDIA H100, B200) that offer superior performance per watt and specialized LLM acceleration features.

Operational Challenges and Engineering Trade-offs

Running LLM inference at Netflix's scale presents unique operational hurdles and requires careful navigation of engineering trade-offs.

- **GPU Resource Fragmentation:** Different LLMs require varying amounts of GPU memory and compute. Efficiently packing diverse models onto a finite set of GPUs without leaving significant idle capacity is a constant challenge.
 - **Trade-off:** Dedicated GPU instances per model (simpler, less efficient) vs. multi-model serving on shared GPUs (complex, higher efficiency). Netflix leans towards multi-model serving with robust resource isolation.
- **Cold Start Latency:** When a new model version is deployed or an underutilized instance scales up, loading large model weights from storage to GPU memory can introduce significant latency.
 - **Solution:** Proactive model pre-loading, instance warm-up strategies, and efficient model streaming.
- **Model Obsolescence and Drift:** LLMs can quickly become outdated, and their performance can degrade over time due to shifts in data distribution or user behavior.
 - **Solution:** Automated retraining pipelines, continuous monitoring of model quality metrics (e.g., perplexity, response relevance), and a robust A/B testing framework.
- **Cost Management:** GPUs are expensive. Balancing performance targets with cost constraints is a continuous effort.
 - **Trade-off:** On-demand instances (flexibility) vs. reserved instances/spot market (cost savings). Netflix uses a hybrid approach, leveraging spot instances for non-critical or batch inference.

- **Observability Granularity:** Monitoring a traditional microservice is simpler than monitoring an LLM inference server. Key metrics include token generation rate, KV cache hit ratio, specific layer latencies, and total cost per token.
 - **Solution:** Custom metrics collection at the inference server level, integrated with Netflix's Atlas monitoring system.

Ensuring Reliability, Observability, and Security

Robust operational practices are non-negotiable for a critical service like LLM inference.

Reliability

- **Regional Redundancy:** Deploying inference clusters across multiple AWS regions and availability zones to withstand regional outages.
- **Circuit Breakers and Rate Limiting:** Implementing these patterns at the API Gateway and Inference Router layers to prevent cascading failures and protect downstream services from overload.
- **Graceful Degradation:** Designing applications to fall back to simpler models, cached responses, or alternative logic if the primary LLM inference service experiences issues.
- **Automated Health Checks:** Continuous monitoring of inference server health, GPU status, and model readiness, with automatic remediation (e.g., restarting unhealthy instances).

Observability

- **Comprehensive Metrics:**
 - **Inference Metrics:** Request QPS, average latency, p99 latency, error rates, token generation rate (tokens/sec), KV cache utilization.
 - **GPU Metrics:** GPU utilization, memory usage, temperature, power consumption.
 - **Business Metrics:** Cost per inference, cost per generated token, model quality scores.
- **Distributed Tracing:** Leveraging OpenTelemetry to trace requests across multiple services, providing full visibility into the path of an inference request and pinpointing bottlenecks.

- **Structured Logging:** Centralized logging with detailed context for every inference request, enabling debugging and post-mortem analysis.
- **Alerting and Dashboards:** Real-time alerts for deviations from performance SLAs or resource thresholds, visualized through custom dashboards built on Netflix Atlas.

Security

- **Data Isolation:** Strict network segmentation ensures that sensitive training data and user prompts do not inadvertently leak between different LLM applications or environments.
- **Access Control:** Fine-grained IAM policies for accessing model artifacts in the Model Store and for invoking inference endpoints.
- **Secure Communication:** All internal and external API calls are encrypted using TLS.
- **Model Integrity:** Mechanisms to verify the integrity and authenticity of model weights before loading them into memory, preventing tampering.

Impact and Future Directions

Netflix's investment in in-house LLM serving has had a profound impact across the organization:

- **Enhanced Personalization:** More nuanced recommendations, personalized content summaries, and tailored search results directly translate to improved user engagement and retention.
- **Accelerated Content Creation:** LLMs assist scriptwriters, localizers, and marketing teams with idea generation, translation, and content analysis, significantly speeding up production workflows.
- **Operational Efficiency:** Automation of tasks like content tagging, metadata generation, and customer support triage leads to substantial operational cost savings.
- **Strategic Independence:** Reduced reliance on external vendors for core AI capabilities, providing greater control over innovation, security, and cost.

Looking ahead, Netflix is exploring:

- **Multi-Modal LLMs:** Integrating vision and audio inputs with LLMs to create even richer, more interactive user experiences.

- **Edge Inference:** Investigating techniques to run smaller, highly optimized models closer to the user (e.g., on smart TVs or mobile devices) for ultra-low latency scenarios.
- **Next-Generation Hardware:** Continuous evaluation of emerging AI accelerators and specialized chips to further improve performance and cost efficiency.

Practical Takeaways for Platform Engineers and AI Developers

Implementing an in-house LLM serving platform is a significant undertaking. Based on Netflix's experience, here are key lessons for others:

1. **Define Clear Requirements Early:** Before writing any code, thoroughly document your latency, throughput, cost, and security requirements. These will drive every architectural decision.
2. **Leverage Existing Infrastructure:** Integrate with your current cloud provider and internal orchestration (e.g., Kubernetes). Don't reinvent the wheel for core services like API gateways, monitoring, or logging.
3. **Invest Heavily in Observability:** LLMs are black boxes. Without deep metrics on GPU usage, token generation, KV cache, and latency, debugging performance issues or cost overruns becomes impossible.
4. **Prioritize Cost Efficiency from Day One:** GPU costs are the primary driver. Implement dynamic batching, quantization, and smart scaling policies from the start. Explore spot instances and reserved capacity.
5. **Embrace Specialized Serving Runtimes:** Generic web servers are insufficient for LLMs. Adopt or build upon frameworks like vLLM or Triton Inference Server that are designed for LLM-specific optimizations.
6. **Build for Heterogeneity:** Anticipate needing to serve multiple LLM architectures and sizes. Design your orchestration and serving layers to be flexible and extensible.
7. **Automate Everything Possible:** From model deployment to scaling and health checks, automation reduces operational overhead and human error.
8. **Security and Data Privacy are Paramount:** Treat LLM inputs and outputs with the same rigor as any sensitive production data. Implement robust access controls and data isolation.

9. **Foster Collaboration:** Close collaboration between MLOps, infrastructure, and AI research teams is essential for success. The interplay between model design and serving infrastructure is critical.
-

References

- Netflix Technology Blog (General MLOps and Infrastructure articles)
 - AWS Machine Learning Blog (For general cloud ML infrastructure patterns)
 - Kubernetes documentation (For container orchestration principles)
 - Open-source project documentation for vLLM, NVIDIA Triton Inference Server, TensorRT-LLM (For specific inference optimization techniques)
 - Academic papers on LLM serving optimization (e.g., continuous batching, FlashAttention)
-

Transparency Note

This case study is a hypothetical reconstruction based on publicly available information regarding Netflix's general engineering practices, industry trends in LLM deployment, and common challenges faced by large-scale AI infrastructure teams. While it aims for technical realism and accuracy, specific implementation details, internal project names, and precise performance metrics are illustrative and not direct disclosures from Netflix.