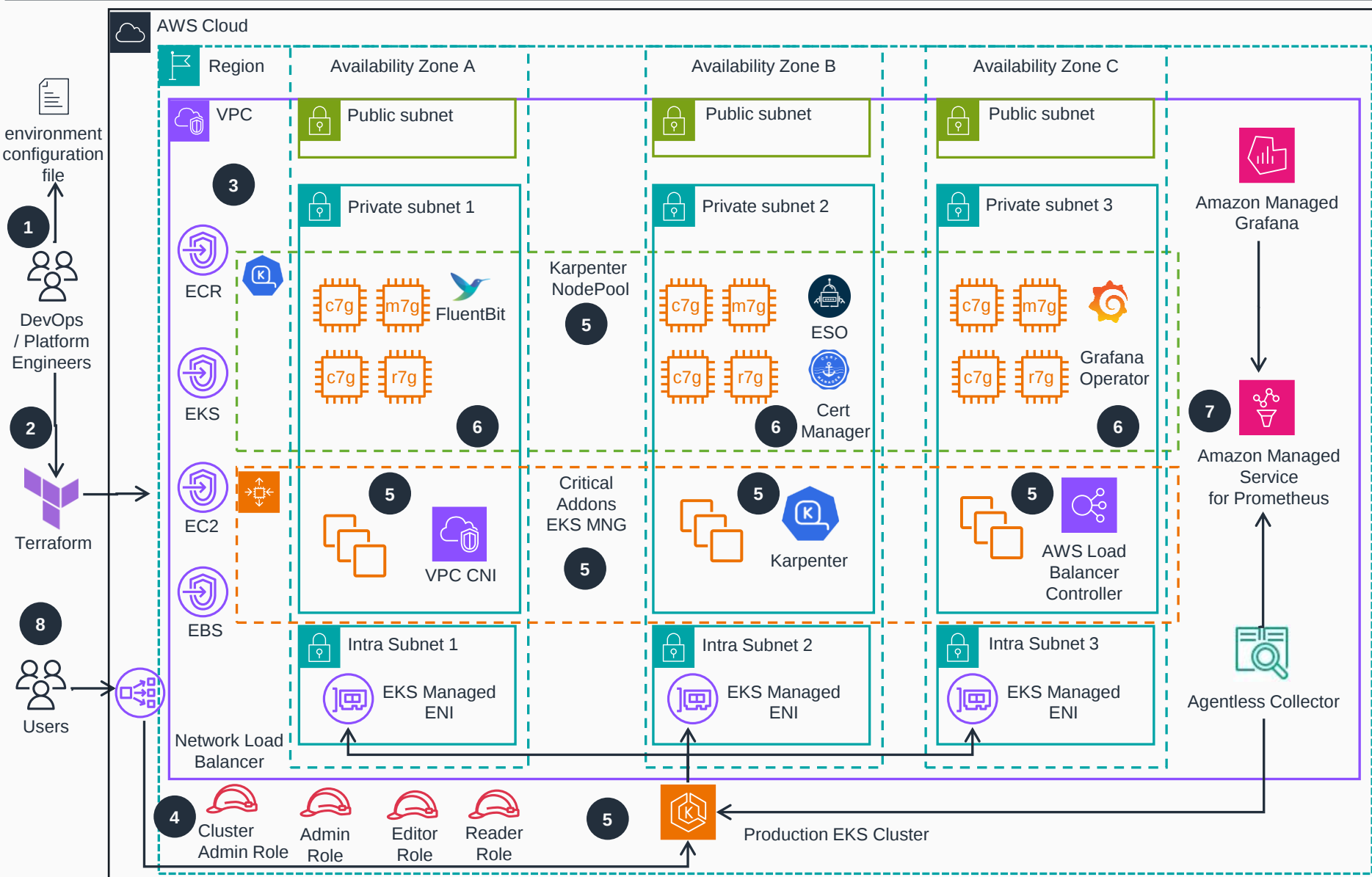


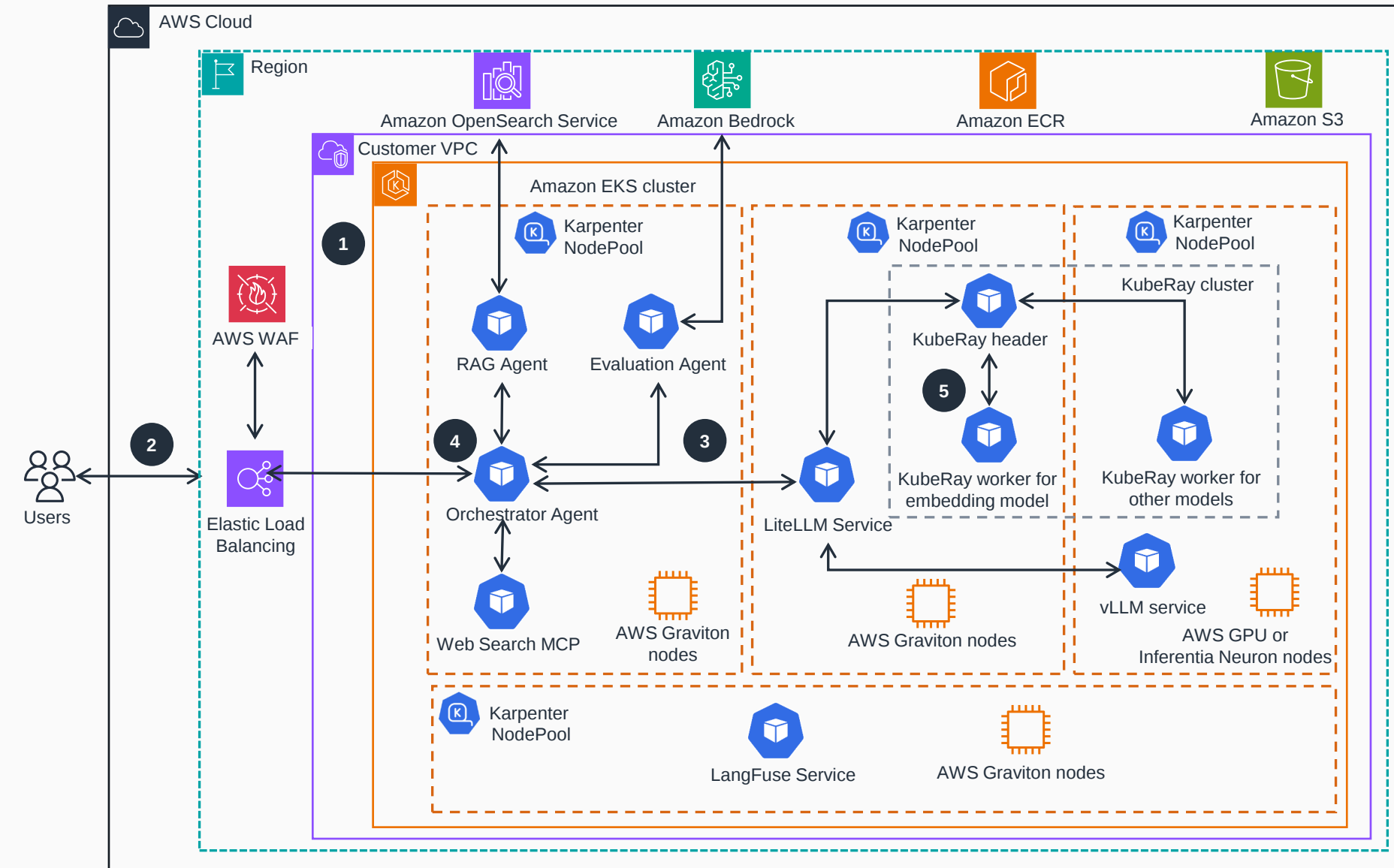
Guidance for Scalable Model Inference and Agentic AI on Amazon EKS

This diagram shows how to provision an Amazon Elastic Kubernetes Service (EKS) cluster with best practices configuration and critical add-ons for AI workloads



Guidance for Scalable Model Inference and Agentic AI on Amazon EKS

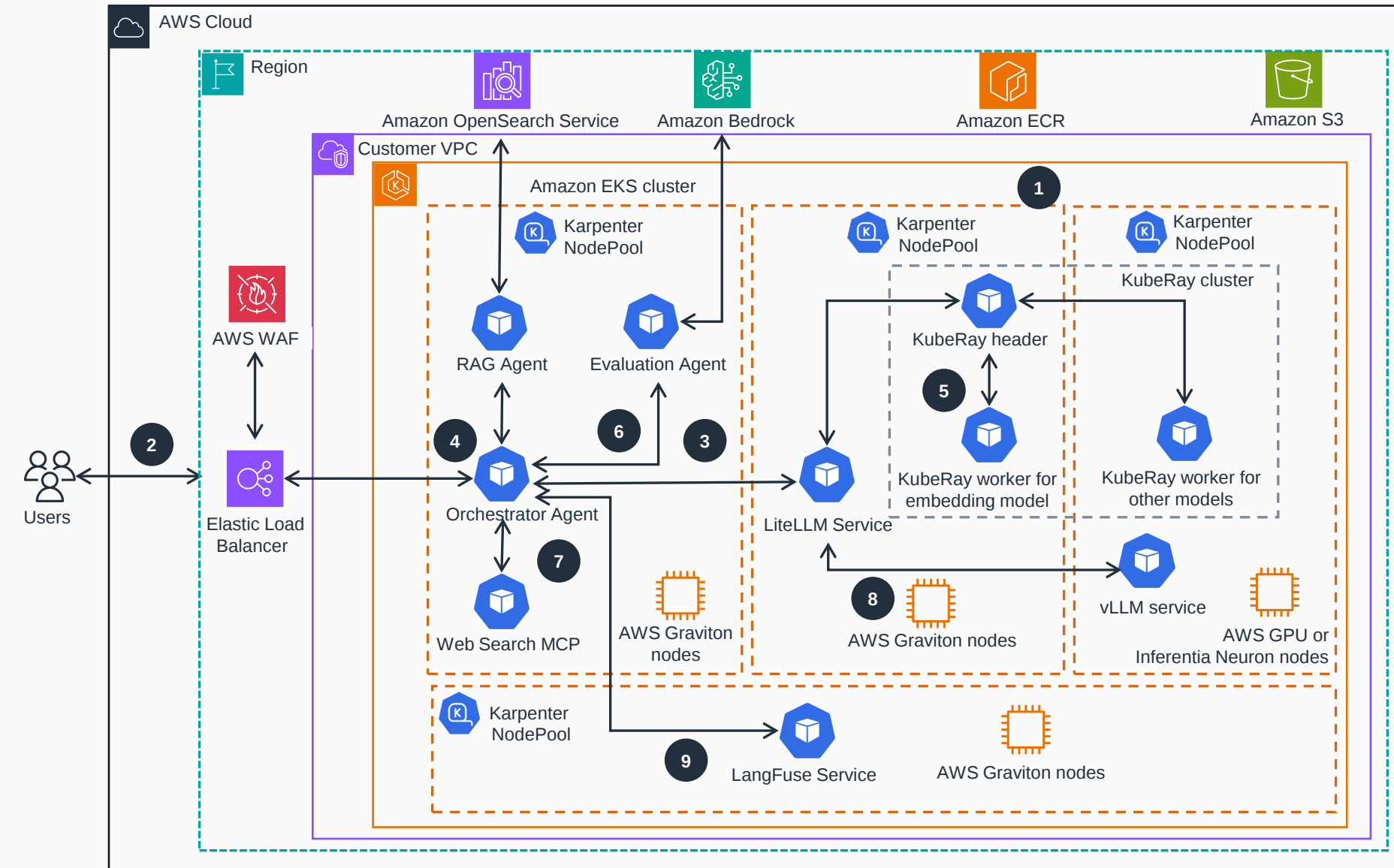
This architecture diagram demonstrates how to deploy ML Models and MCP server and agent on EKS cluster, provide unified model inference API Gateway, and implement Agentic AI capabilities to enable models' interactions.



- Guidance workloads are deployed into an Amazon EKS cluster, configured for application readiness with compute plane managed by Karpenter auto-scaler. This setup efficiently auto-scales AWS Graviton, GPU and AWS Inferentia based compute nodes across multiple Availability Zones (AZs), ensuring robust infrastructure distribution and high availability of various Kubernetes services.
- User interaction starts with an authentication to the EKS cluster via IAM. Application HTTP(s) requests are directed to an exposed endpoint, managed by Elastic Load Balancing and protected by AWS Web Application Firewall (WAF). This endpoint URL serves as the gateway for all user queries and ensures balanced distribution of incoming traffic with consistent accessibility.
- The Orchestrator agent, powered by Strands Agent SDK, serves as the central workflow coordination hub. It processes incoming queries by connecting with reasoning models supported by LiteLLM and vLLM services, analysing the requests, and determining the appropriate workflow and tools needed for response generation. LiteLLM functions as a unified API gateway proxy for model management, providing centralized security controls and standardized access to both embedding and reasoning models.
- Orchestrator agent delegates control to RAG agent to verify knowledge base validity to initiate Knowledge validation. When updates are needed, the RAG agent initiates the process of embedding new information into the Amazon OpenSearch cluster, ensuring the Knowledge Base remains current and comprehensive.
- KubeRay cluster handles the embedding process where the Ray header dynamically manages cluster worker node scaling based on resource demands. These worker nodes execute the embedding process through the llamacpp framework, while the RAG agent simultaneously embeds user questions and searches for relevant information with the OpenSearch cluster.

Guidance for Scalable Model Inference and Agentic AI on Amazon EKS

This architecture diagram demonstrates how to deploy ML Models and MCP server and agent on EKS cluster, provide unified model inference API Gateway, and implement Agentic AI capabilities to enable models' interactions.



1. The Orchestrator agent receives user requests and initiates the process.
2. The Orchestrator agent interacts with the Elastic Load Balancer and AWS WAF.
3. The Orchestrator agent interacts with the RAG Agent.
4. The RAG Agent interacts with the Web Search MCP.
5. The KubeRay header manages the KubeRay workers.
6. The Evaluation agent performs Response Quality assurance which leverages models hosted in Amazon Bedrock. This agent implements a feedback-based correction loop using RAGAS metrics to assess response quality and provides relevancy scores to the orchestrator agent for decision-making purposes.
7. When the RAG agent's responses receive relevancy scores below the configured threshold, the Orchestrator agent initiates a web search process. It retrieves the Tavily web search API tool from the Web search Model context Protocol (MCP) server and performs dynamic queries to obtain supplementary or corrective information.
8. Models based on vLLM framework running on GPU (or Inferentia) EKS compute nodes generate final responses which are returned via LiteLLM gateway to the Orchestrator agent. The reasoning model processes the aggregated information, including both Knowledge Base data and Web search results when applicable, refining and rephrasing the content to create a coherent and accurate response for the user's prompt.
9. Throughout this entire process, the Orchestrator agent maintains comprehensive interaction tracking. It automatically traces all agent activities and communications, with the resulting metrics being visualized via the LangFuse service, providing transparency and monitoring of the system's operations. Additional EKS infrastructure level observability is available via previously deployed and configured Prometheus /Grafana monitoring stack.