

Advanced LLM Concepts

Transformer architecture, attention mechanisms, RAG, embeddings, and scaling laws — a clear technical guide for anyone building with LLMs.

Dr. Farshid Pirahansiah · pirahansiah.com



Mind Map Advanced LLM Concepts

Mind Map: Orchestrating Agents & Advanced LLM Concepts

1. Introduction

- **Main Concept:** Coordination of multiple AI agents for complex tasks, enhanced by Large Language Models (LLMs) and multimodal technologies.
- **Goals:** Solve tasks beyond a single agent's capability, using agents and LLMs in harmony.
- **Key Technologies:** LLMs, multimodal models (vision + text), task orchestration.

2. Core Components of Agent Systems

2.1 Agents

- **Definition:** Autonomous entities carrying out specific functions.
- **Types:**
 - **Single-purpose:** Designed for specific tasks.
 - **General-purpose:** Flexible agents capable of performing various tasks.
- **Capabilities:**
 - **Interaction** with environments.
 - **Processing inputs** and generating outputs.
 - **Self-contained** decision-making.

2.2 Orchestrator

- **Definition:** Central controller managing multiple agents and interacting with LLMs.
- **Roles:**
 - Delegates tasks to agents.
 - Monitors progress.
 - Facilitates communication between agents and LLMs.
 - Combines results for task completion.

2.3 Communication Between Agents and LLMs

- **Importance:** Efficient communication enables effective agent-LLM collaboration.

- **Methods:**
 - **Message Passing** between agents.
 - **API Calls** to trigger LLMs for specific operations (like function calling).
 - **Shared Memory** or databases for data exchange.

3. Advanced LLM Techniques

3.1 Function Calling

- **Definition:** Triggering LLM functions based on tasks, allowing for interaction with external APIs.
- **Similar Concepts:**
 - API Calls.
 - Agent Invocation.

3.2 Subtask Processing

- **Definition:** Dividing complex tasks into smaller subtasks handled by individual agents or LLMs.
- **Similar Concepts:**
 - Task Decomposition.
 - Multi-step Learning.

3.3 Tools Through Prompting

- **Definition:** Prompting LLMs to interact with external tools, such as databases or browsers, enhancing their capability.
- **Similar Concepts:**
 - Tool-Augmented Models.
 - Retrieval-Augmented Generation (RAG).

4. Multimodal LLM Technologies

4.1 VisionLLM v2

- **Overview:** Handles both vision and language tasks. Excels in **image segmentation**, **object detection**, and **caption generation**.
- **Use Case:** Medical image analysis or automated content creation where visual and text-based processing is needed.

4.2 NVLM 1.0 (NVIDIA)

- **Overview:** Specializes in **OCR**, **visual reasoning**, and **localization**. Handles complex vision-text integrations for real-world applications.
- **Use Case:** Autonomous driving, where vehicles and pedestrians are detected in real-time using visual cues.

4.3 Apple's 4M Multimodal Model

- **Overview:** Excels in **image editing** and **3D scene manipulation** based on natural language commands.
- **Use Case:** Advanced photo editing or 3D modeling, e.g., "brighten the sky" or "remove an object" via commands.

4.4 Molmo

- **Overview:** Designed for **dense image captioning** and **object pointing**, using the **PixMo** dataset.
- **Use Case:** Geospatial analysis, where satellite images can be analyzed and annotated for buildings or geographic markers.

4.5 Qwen2-VL

- **Overview:** Multimodal model that integrates **video comprehension** and **multilingual text recognition**.
- **Use Case:** Sports analytics, tracking player performance and answering questions related to game footage.

4.6 Pixtral by Mistral

- **Overview:** Capable of **multi-image processing** at native resolutions with an outstanding ability to follow instructions.
- **Use Case:** Security monitoring by processing multiple camera feeds for detecting and tracking objects in different locations.

5. Embedding and Vector Technologies

5.1 Knowledge Graphs

- **Definition:** Represents structured knowledge by entities and their relationships to enhance reasoning in LLMs.
- **Use Cases:** Question-answering, data retrieval.

5.2 pgvector

- **Definition:** PostgreSQL extension for handling vectors in similarity search and embedding-based queries.
- **Use Cases:** Efficient storage and search of embeddings in recommendation systems or LLM-driven queries.

6. Localization of LLMs

- **Definition:** Optimizing LLMs for region-specific tasks, addressing local language nuances and data regulations.
- **Similar Concepts:**
 - Language Adaptation.
 - Regional Compliance.

7. Advanced Prompting and RAG Techniques

7.1 Superposition Prompting

- **Definition:** Layering multiple prompts for richer context, improving the LLM's decision-making.
- **Use Cases:** Multi-step decision-making tasks.

7.2 RAG in Computer Vision (CV)

- **Definition:** Combines retrieval systems with generation models, leveraging external image data for visual task processing.
- **Use Case:** Document processing where both text and images need to be integrated, such as legal document review.

7.3 Embedding Models

- **Definition:** Vectorizing input data (images, text) for similarity search or clustering in LLM workflows.
- **Use Cases:** Product recommendations, clustering similar documents for faster access.

8. Productionizing LLMs

8.1 Production-Ready LLMs

- **Definition:** Deploying LLMs in real-world applications, considering performance, cost, and scalability.
- **Use Cases:** Chatbots, real-time language translation.

8.2 LLM with RAG (Retrieval-Augmented Generation)

- **Definition:** Enhancing LLMs by coupling them with retrieval-based systems for information generation, albeit at higher costs.
- **Use Cases:** Advanced customer service, legal text generation from large databases.

8.3 LLM with IoT

- **Definition:** Integrating LLMs with IoT devices to process data and control smart systems through natural language interfaces.
- **Use Cases:** Home automation, real-time smart sensor monitoring.

9. Example Workflow

9.1 Research Task

- **Step 1:** Task is divided into subtasks (research, summarizing, final report).
- **Step 2:** Agents handle subtasks, and LLMs assist with content generation and summarization.
- **Step 3:** Results from agents are aggregated into a cohesive report by the orchestrator.

9.2 Multi-Agent & LLM Collaboration

- **Scenario:** Writing a complex essay.

- **Agent 1:** Researches the topic.
- **Agent 2:** Summarizes key points.
- **Agent 3:** Drafts the essay.
- **LLM:** Enhances writing style and content.
- **Orchestrator:** Oversees, revises, and finalizes the report.

10. Benefits

- **Efficiency:** Tasks are completed faster with parallel agent and LLM collaboration.
- **Scalability:** Able to handle larger, more complex tasks with both agents and LLMs.
- **Flexibility:** Agents and LLMs can adapt to different scenarios.
- **Improved Decision-Making:** Real-time adaptability via dynamic task handling.

11. Challenges

- **Coordination Complexity:** Synchronizing multiple agents and LLMs.
- **Communication Overhead:** Managing the communication between agents and LLMs.
- **Error Handling:** Failures can affect the entire workflow.
- **Cost & Resources:** High costs, especially with retrieval and multimodal systems like RAG.

12. Applications

- **Research & Analysis:** For deep analysis and data extraction across domains.
 - **Content Creation:** Multi-agent systems producing high-quality written content.
 - **Project Management:** Automating and organizing complex, large-scale projects.
-

- Function Calling
 - Subtasks
 - Tools Using Prompting
 - Knowledge Graphs
 - pgVector
 - open source LLM
 - Superposition Prompting
 - Retrieval-Augmented Generation (RAG) in Computer vision
 - Embedding Models
 - Production-Level LLMs
 - Production-Level LLMs with RAG
 - Production-Level LLMs in IoT
 - Cost Management
-

- In managing costs for production-level LLMs with RAG, strategies like embedding retrieval systems and optimizing model queries can help reduce resource use. Leveraging cloud-based solutions or fine-tuning smaller models for specific tasks are also common methods to reduce the expense of running these models in real-time .
- Running LLMs in production at scale can be expensive, particularly when combined with RAG (Retrieval-Augmented Generation). This is because both the LLM and the external retrieval system (like a knowledge graph) need to work in tandem, which can increase computational overhead. When applied to IoT (Internet of Things), where data sources are vast and diverse, the cost of maintaining such systems can become prohibitive .
- Embedding models help LLMs convert text into vector representations, enabling tasks like search, similarity comparison, and context retrieval. These models, when integrated with LLMs, allow for scaling up chunk processing for large datasets, making them essential for production environments where real-time data retrieval is required .
- Function calling within LLMs allows these models to interact with external APIs and tools by converting user queries into function calls. For example, if a user asks about the weather, the model can trigger a function like `get_current_weather(location)`, passing the necessary parameters to retrieve the correct data. This capability enhances the integration of LLMs with external systems, enabling a more interactive and task-specific approach .
- When combined with subtasks, LLMs can divide complex tasks into manageable parts. For instance, a user query can trigger multiple functions or tools to handle different subtasks, such as searching databases or summarizing documents . Tool integration with LLMs through prompting enables the models to call upon specialized tools for tasks that require external data or specific functionality. For example, using tools like LangChain, LLMs can be prompted to interact with knowledge graphs, retrieve data from external APIs, or handle subtasks like document chunking or embedding generation . These prompts allow the LLM to understand when and how to invoke external resources, ensuring more accurate responses.
- A knowledge graph allows LLMs to structure unstructured data by mapping entities and relationships in a graph format, making it easier for models to retrieve, reason about, and respond to queries based on complex datasets . Tools like pgVector enhance this by enabling vectorized searches, allowing models to retrieve semantically similar content from large datasets efficiently.
- Superposition prompting is an advanced technique that presents multiple prompts to the model at once, helping it learn or decide between different potential tasks. This method increases the model's ability to handle ambiguity or multiple concurrent tasks efficiently. The application of such techniques can be seen in sectors like Apple, where AI-driven interfaces need to handle complex decision-making processes.
- Keyword Search vs. Vector Search:
 - When comparing keyword search and vector search, an alpha-based selection helps determine which method is more appropriate depending on the data and task.
 - Chunking Data: Data can be divided or "chunked" into meaningful units such as paraphrases, pages, or chapters. This approach helps structure information for more efficient retrieval and processing.
 - Imagine an 8MP

picture: Whether it's an image of one person or many people, the resolution remains 8MP. In the same way, the numerical representation of data, like vectors, maintains its "size" regardless of what is being represented. • GraphQL: GraphQL offers a flexible query system, enabling precise data retrieval by requesting only the fields needed, making it an effective alternative to traditional REST APIs in many cases.

- Best Methods for Keyword Search:

- For a keyword-only search, BM25 is considered one of the most effective models. It excels in ranking documents based on keyword occurrences and relevance.
 - To measure keyword frequency or how common a keyword is, approaches like term frequency (TF) and inverse document frequency (IDF) can provide insight into the importance of keywords within a dataset.