

Introduction to RAG

Retrieval-Augmented Generation

- **Definition**

Technique that enhances LLMs by retrieving external, private data before generating a response.

- **Goal**

To ground AI in factual, company-specific information.

- **Bypasses Knowledge Cut-offs**

Accesses the most current information by querying real-time data sources.

- **Reduced Hallucinations**

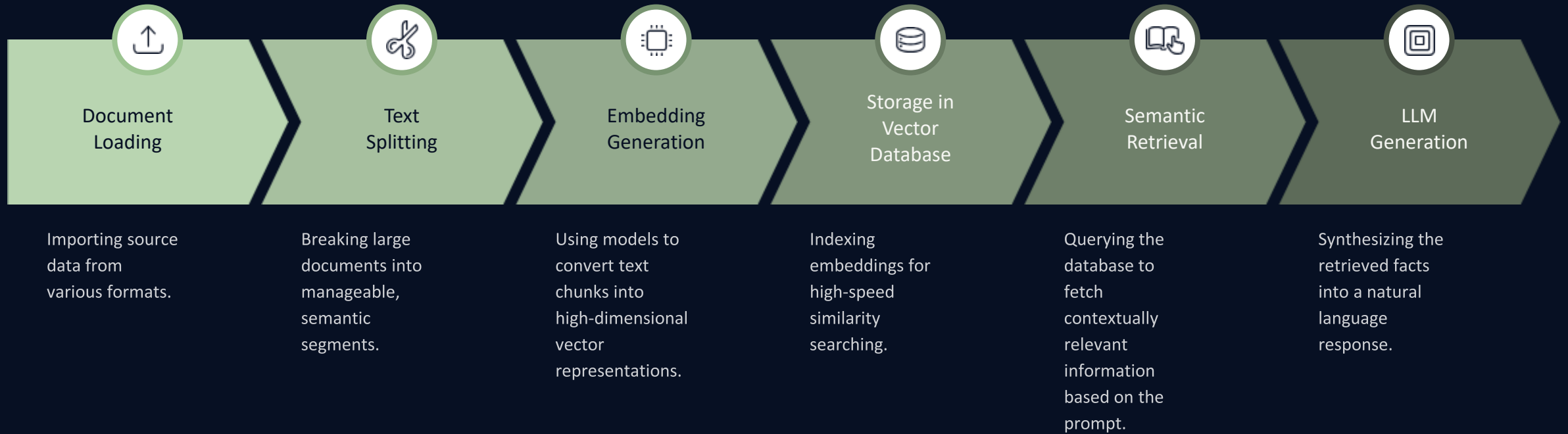
Provides relevant context, forcing the model to base answers on provided documentation.

- **Privacy & Compliance**

Ensures AI interactions rely on internal documents, allowing for granular access control.

The RAG Architecture Pipeline

A step-by-step breakdown of the Retrieval-Augmented Generation process



Document Processing & Chunking

Efficient strategies for data ingestion and semantic segmentation



Document Loading

Efficient ingestion of diverse file formats including PDFs, Word documents, and HTML.



Recursive Character Splitting

Iteratively splitting text while prioritizing the preservation of sentence and paragraph structure.



Text Splitting

Breaking down large, monolithic files into smaller, manageable, and highly contextual pieces.



Semantic Chunking

Grouping text based on conceptual meaning to ensure high-quality context retention for better retrieval accuracy.



Fixed-size Chunking

Segmenting text into uniform blocks based on a character count (e.g., 500 characters).

Embeddings & Vector Databases

Transforming data into intelligent insights

- **Embeddings**

Converts text into high-dimensional numerical vectors to capture deep semantic meaning and relationships between data points.

- **Vector Databases**

Specialized storage engines like FAISS, Chroma, and Pinecone, optimized for high-speed similarity searching and index management.

- **Semantic Search**

Enables retrieval based on conceptual meaning rather than literal keywords, significantly improving relevance by understanding user intent.

Advanced Retrieval Techniques

Optimizing search performance and accuracy



- **Hybrid Search**

Combines semantic search with keyword (BM25) search to maximize precision and recall.

- **Metadata Filtering**

Uses tags, dates, and categorical data to narrow result sets and improve search efficiency.

- **Reranking**

Utilizes cross-encoder models to perform deeper analysis and re-order top retrieval results by relevance.

- **Query Rewriting**

Transforms ambiguous user questions into structured, optimized queries to improve retrieval outcomes.

Advanced RAG Frameworks

Modern architectures and evaluation methodologies for retrieval-augmented systems



Agentic RAG

Utilizes autonomous agents to dynamically choose between information retrieval and reasoning based on query complexity.



Graph RAG

Integrates Knowledge Graphs to map entity relationships, facilitating superior context for complex, multi-hop queries.



RAG Evaluation

Uses frameworks like RAGAS to quantify performance via metrics: Faithfulness, Answer Relevance, and Context Precision.

Tools & Ecosystem

Key frameworks and vector database solutions

- **LangChain**

Framework for orchestration and workflow management.

- **LlamaIndex**

Specialized tool for data indexing and retrieval ingestion.

- **Pinecone**

Cloud-managed, scalable vector storage solution.

- **Chroma**

Open-source, developer-friendly local database.

- **Qdrant**

High-performance engine for advanced retrieval tasks.

- **FAISS**

Efficient library for local similarity search operations.

Major Project: PDF Knowledge Assistant

End-to-end RAG assistant development roadmap



PDF Ingestion

Extract text and data from technical documentation.



Data Processing

Implement robust chunking strategies and generate vector embeddings.



Model Integration

Utilize high-performance LLMs like GPT-4 or Llama 3 for reasoning.



UI Development

Design an intuitive interface using Streamlit or Chainlit.



Evaluation

Monitor and validate response accuracy to minimize hallucinations.