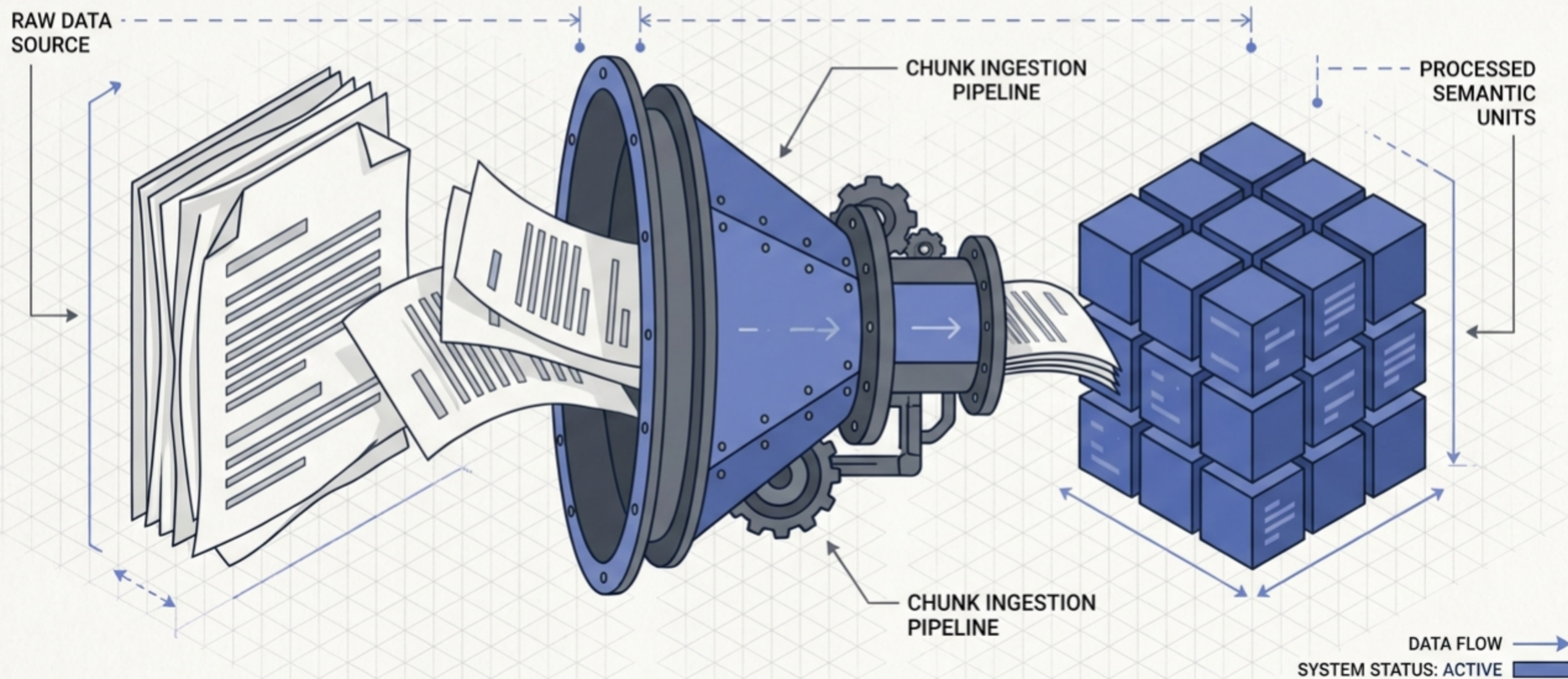


Search Semantically: Chunking & Production Pipelines

Building the ingestion engine that makes Retrieval-Augmented Generation (RAG) real.

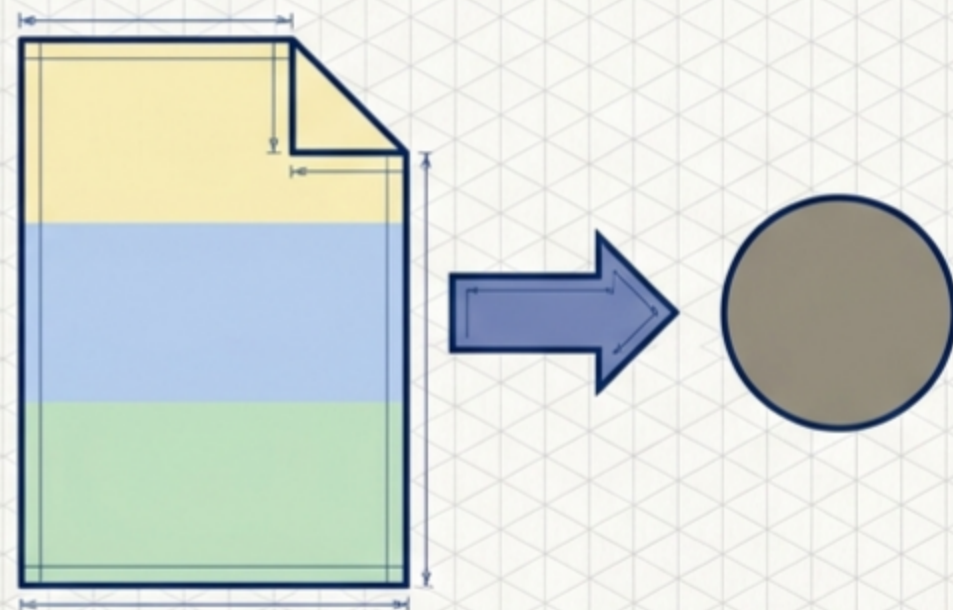


The Context Bottleneck

Without chunking, large documents break RAG architectures in distinct ways. We must isolate the relevant paragraph, not the entire file.



Context Window Overflow



Retrieval Imprecision

One vector for a long document blurs multiple distinct topics into a useless average.

The Mechanics of the Cut: Preserving Context with Overlap

Cutting a document into fixed sizes inevitably slices through connected thoughts. Overlap creates a safety net so no meaning is lost at the boundary.

Without Overlap

The financial results were primarily driven by an unexpected surge in Q4 enterprise sales.



With Overlap

The financial results were primarily driven by an unexpected surge in Q4 enterprise sales.

Chunk A

The Straddled Thought:
Preserved.

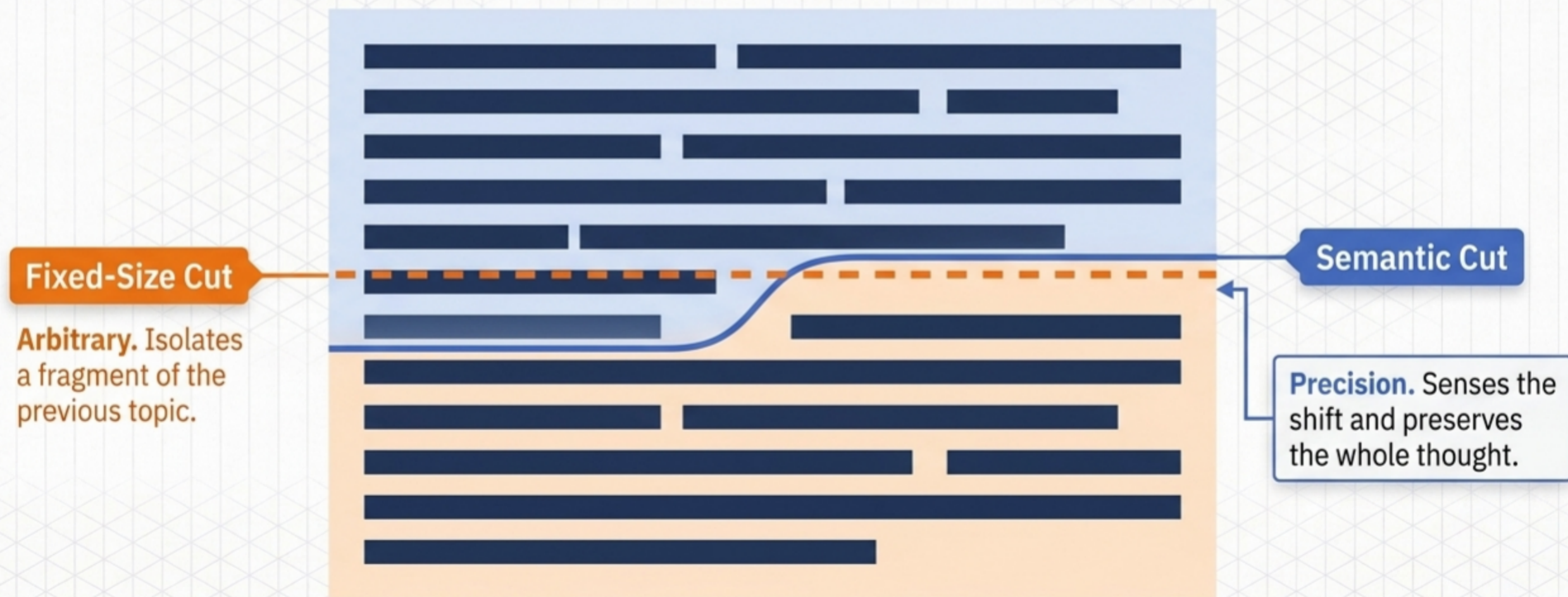
Chunk B

Evaluating Chunking Strategies

	Fixed-Size 	Sentence {,} 	Semantic 
How it cuts	Every N words with an O-word overlap.	Groups whole, complete sentences.	Splits dynamically where the underlying topic shifts.
Pros	Simplest to implement. Predictable processing.	Reads naturally to the LLM. Respects basic grammar.	Smartest grouping. Highly precise retrieval.
Cons	Cuts mid-thought. Arbitrary boundaries.	Varying chunk sizes. Can separate highly related sentences.	Most complex to build. Computationally heavier.

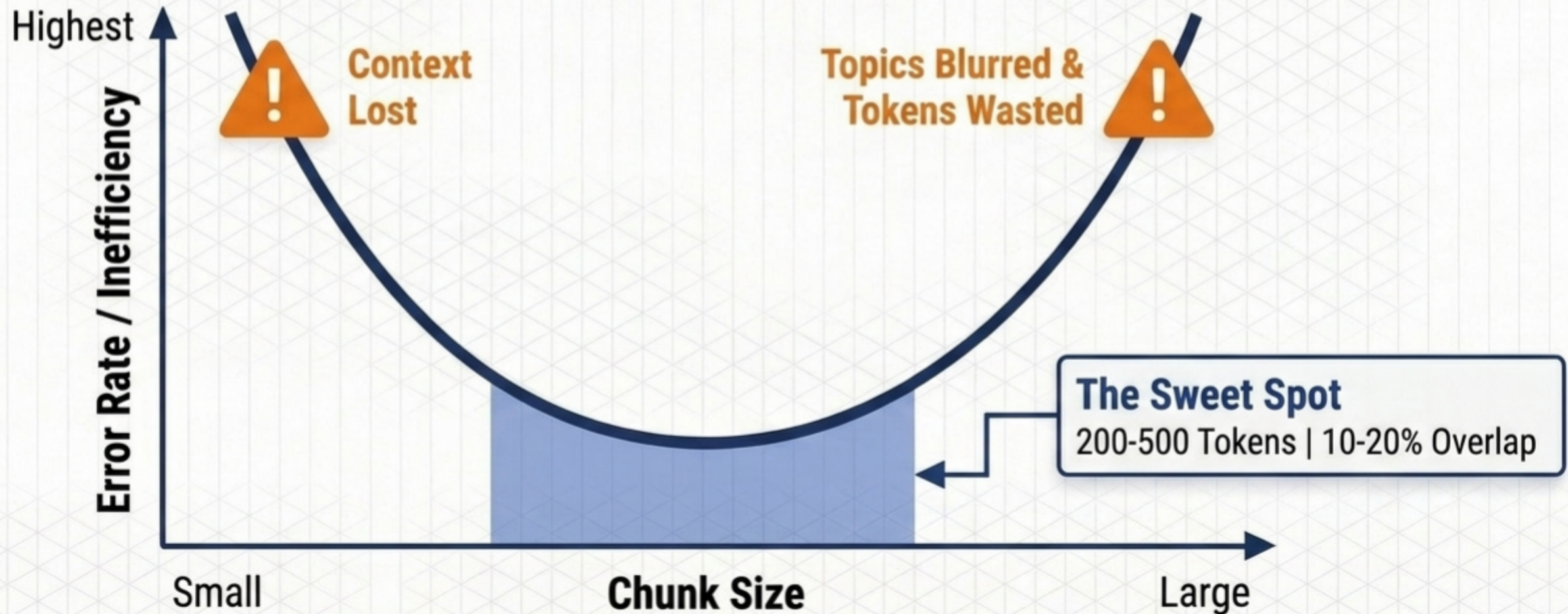
Deep Dive: The Precision of Semantic Chunking

Unlike fixed rules, semantic chunking analyzes meaning to find natural breakpoints in the narrative.



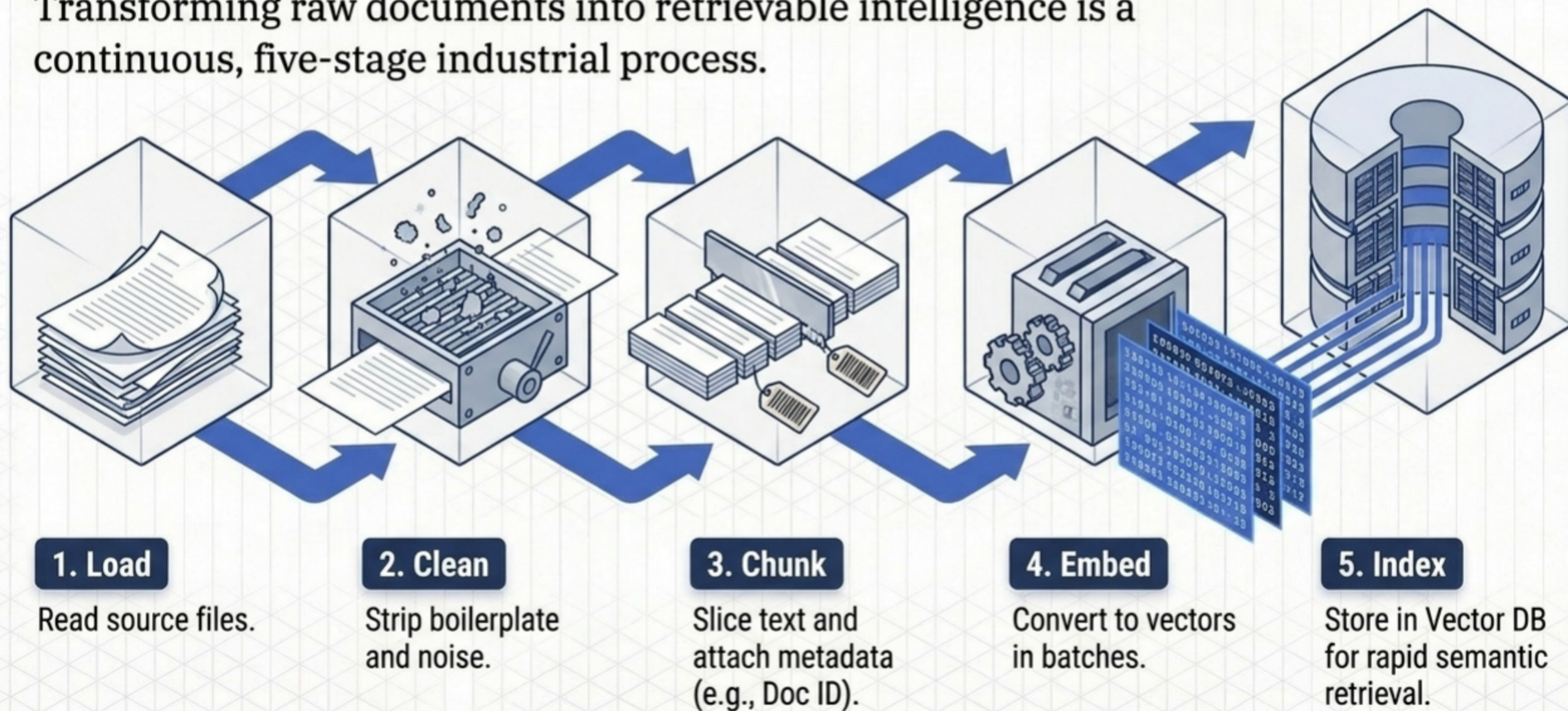
Calibrating the Pipeline: The Goldilocks Size

Tuning chunk parameters requires balancing context retention against precision and cost. Start with the baseline, then tune by measuring RAG quality.



The RAG Ingestion Pipeline

Transforming raw documents into retrievable intelligence is a continuous, five-stage industrial process.

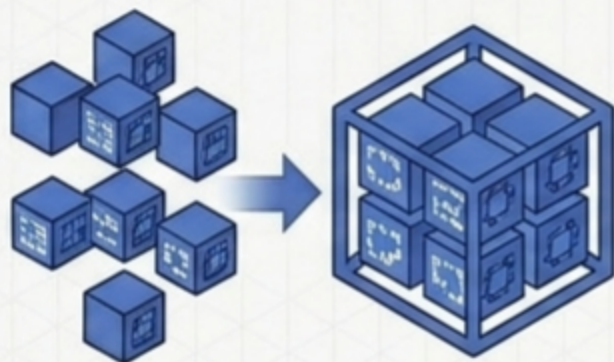


Operating at Scale

A production pipeline requires constant operational oversight to maintain freshness, speed, and accuracy.

Efficiency

Batching
Embeddings to save
compute time.



Cost Control

Caching unchanged
documents to save
processing compute.



Freshness

Planning re-indexing
to ensure up-to-date
answers.



Monitoring



Cost

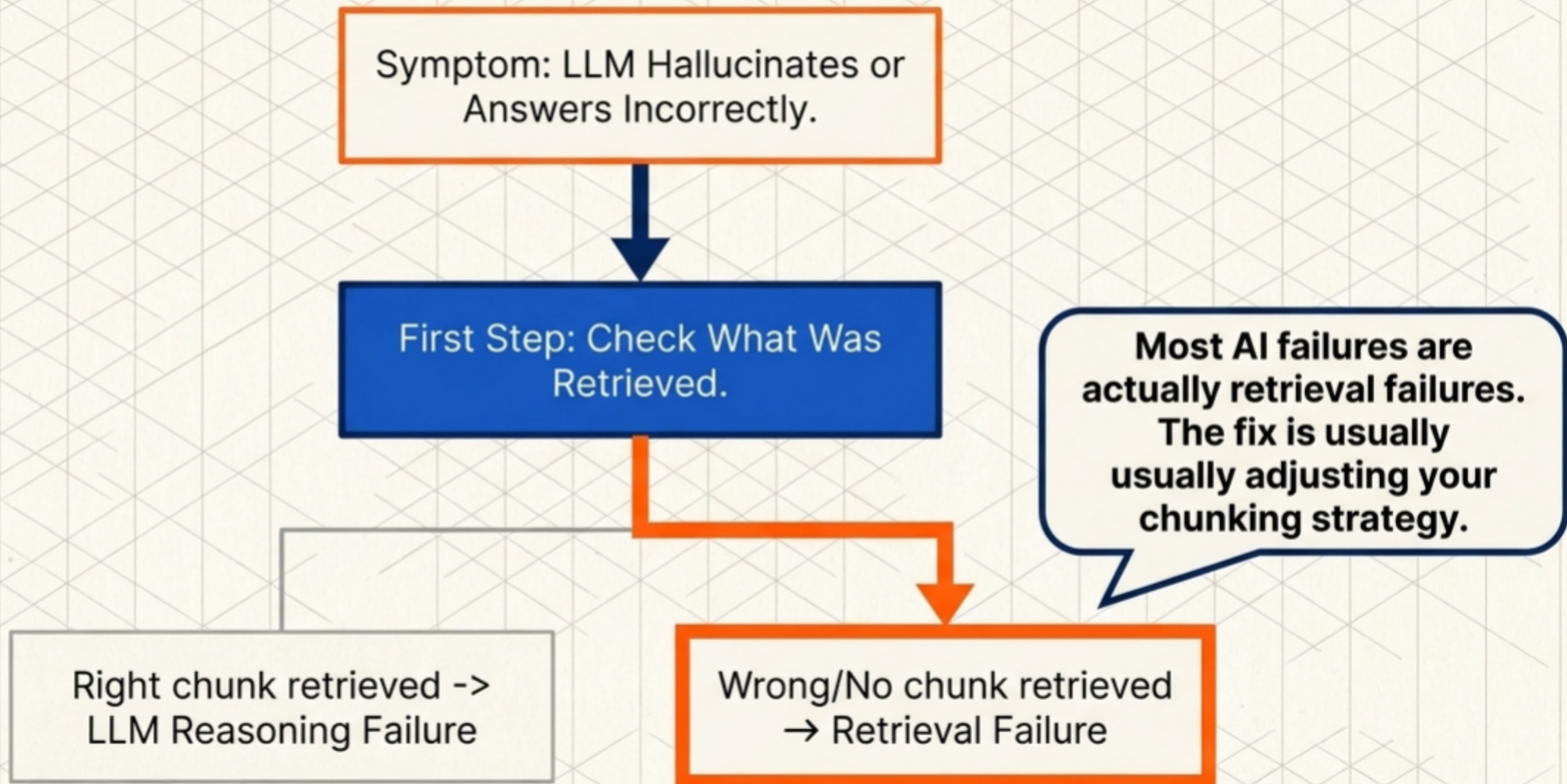


Latency

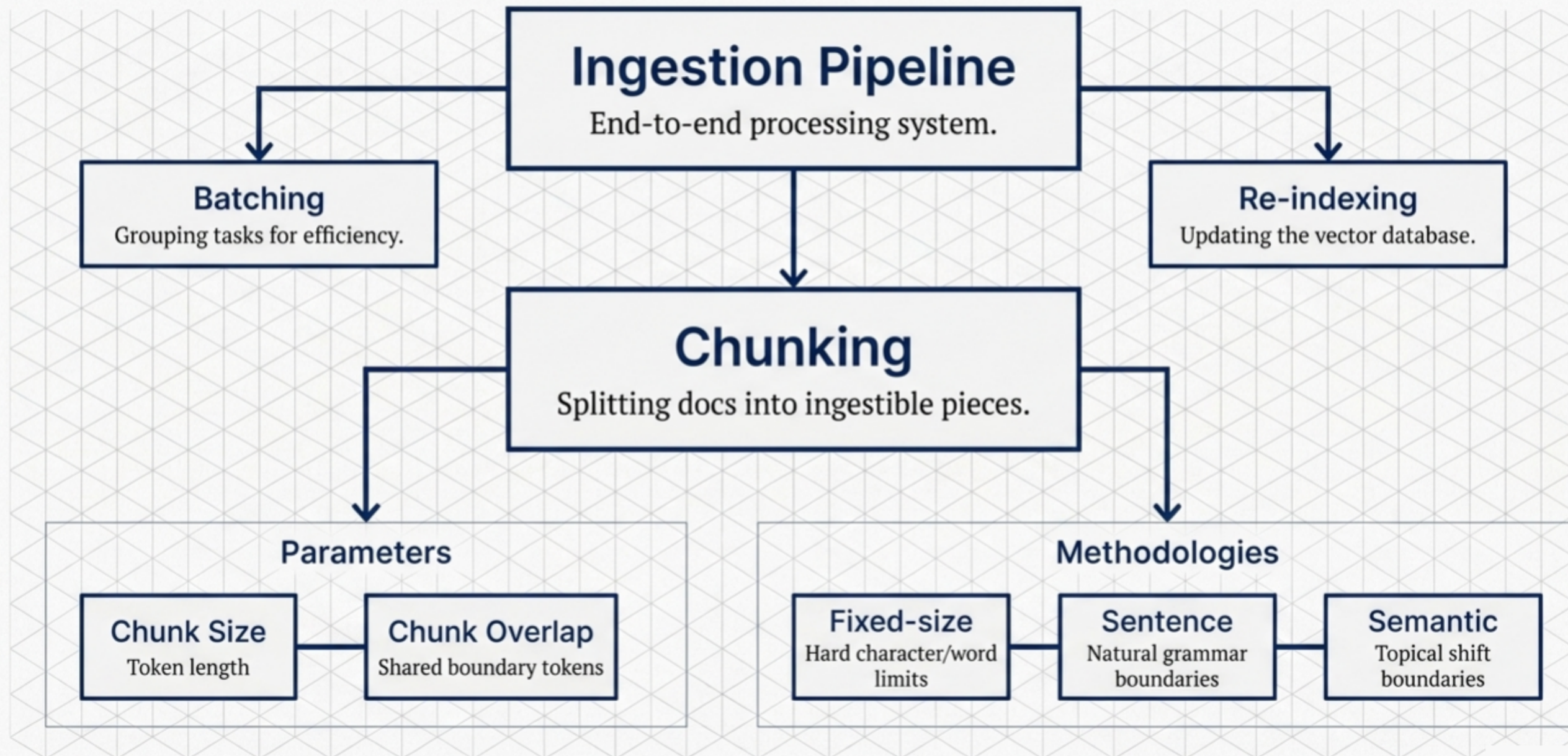


Quality

Quality Control: Debugging RAG



The Blueprint Glossary



Check Your Understanding

The core principles of semantic search ingestion.

Q: Three reasons to chunk long documents?

A:

1. Overcome model input limits.
2. Improve retrieval precision.
3. Provide the LLM with only the exact relevant paragraph.

Q: What critical problem does *overlap* solve?

A: It stops a single thought that straddles a chunking boundary from being destroyed or lost.

Q: A RAG answer is wrong—what do you check first?

A: Whether the right chunk was retrieved.

Most failures are ingestion/retrieval failures, not LLM reasoning failures.