

Generative AI in supervision

The ChatDNB case

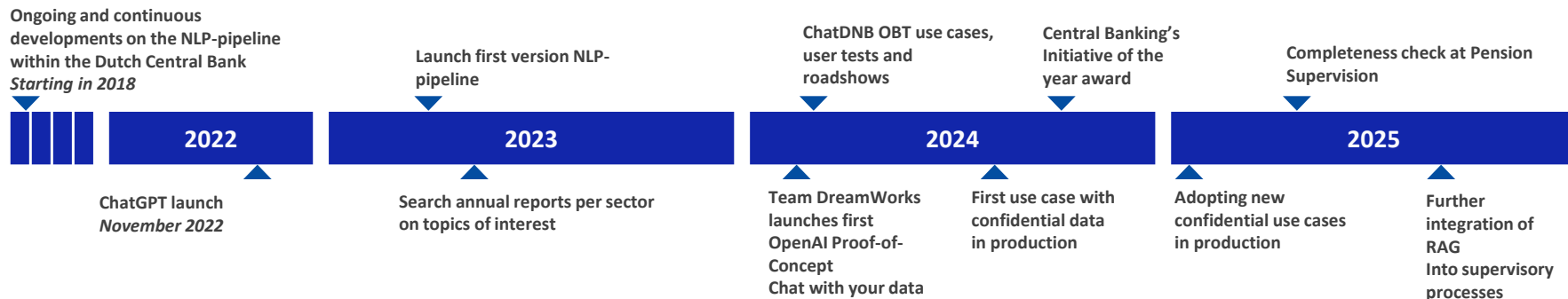
Luc Severeijns & Luca Mossink

DeNederlandscheBank

EUROSYSTEM

Due to years of investment in NLP techniques, DNB was able to scale up quickly after the launch of ChatGPT at the end of 2022, eventually leading to ChatDNB

- Within DNB, Team DreamWorks was already focusing on NLP¹ to search through large amounts of texts automatically
- Following the launch of ChatGPT, Microsoft conducted a demo for our organization to showcase the capabilities of GenAI
- This sparked numerous experiments and discussions about the use of LLMs² within the organization
- A public data use case was chosen to learn from and adopt this new technology before applying it to other scenarios
- After extensive discussions on data security, a first confidential data use case was identified and brought to production
- As colleagues became enthusiastic about the possibilities of LLMs, more use cases were onboarded



ChatDNB makes it possible to service supervisors with GenAI using their data and keeping it a modular solution



Securely chat with data classified up to **DNB-Confidential**



Automatically assign access rights ensuring you only see data you are authorized to view



Search through SharePoint files using an automated integration that retrieves documents and metadata



Filter documents or folders similarly to SharePoint, allowing you to focus on data relevant to your query



Mitigate hallucination by referencing relevant sources and displaying the thought process



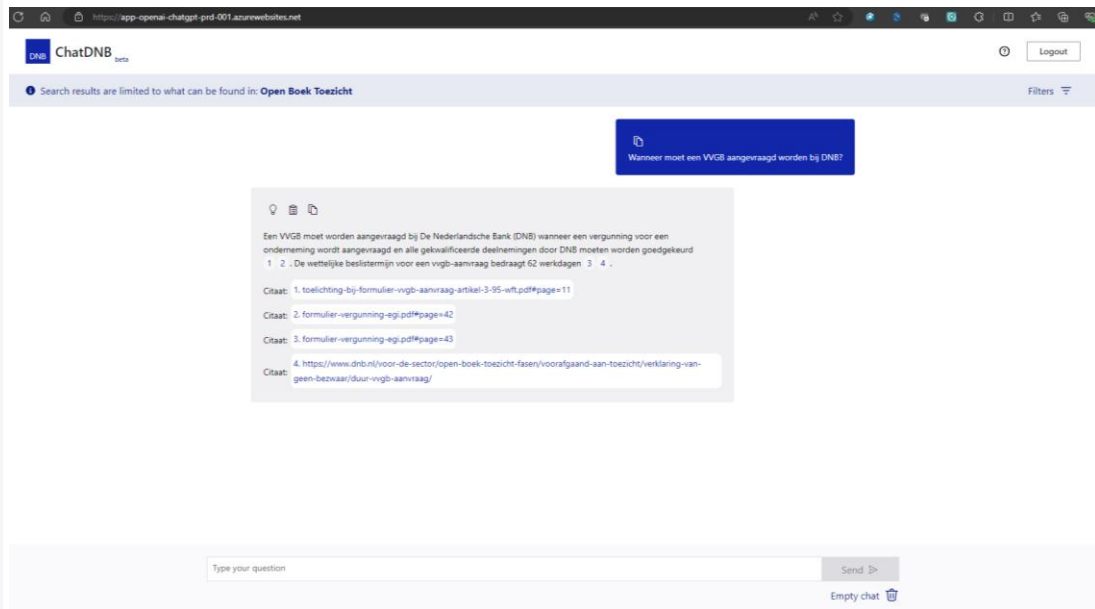
Integrate with supervisory processes to speed up throughput and remain agile for technological changes



Provide feedback to improve the quality of ChatDNB's answers



Monitor the overall quality of generated answers to improve the pipeline and build user trust



There are five categories of use cases that we serve, all powered by a form of RAG¹



Search your dossier

Find relevant information faster with to ad-hoc questions using advanced filter options



Generate draft texts

Receive drafts of texts based on selected relevant documentation/answers and example texts



Completeness Check

Give an overview of the available information for a given topic to assess whether sufficient information has been shared



Consistency Check

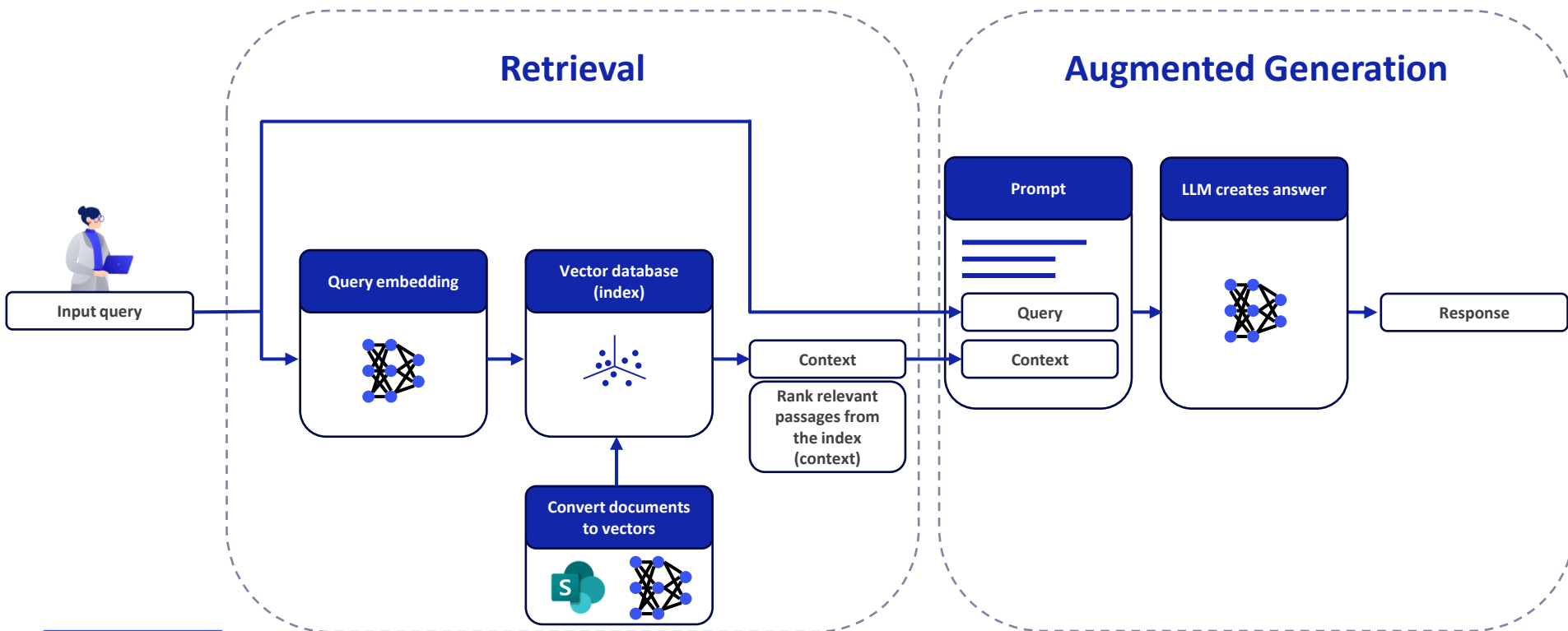
Compare the answers from a questionnaire to the answers given in the text



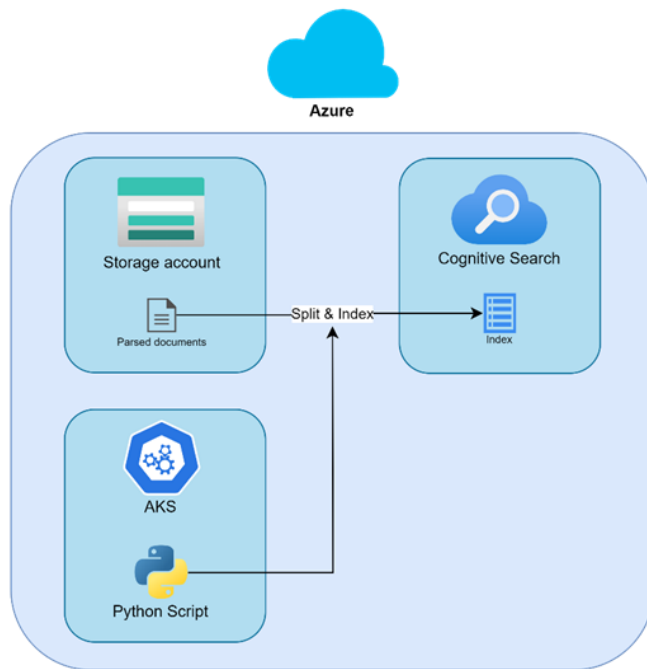
Pre-fill Assessment Frameworks

Get an evidence-based suggestion for an assessment considering relevant regulations and/or guidelines

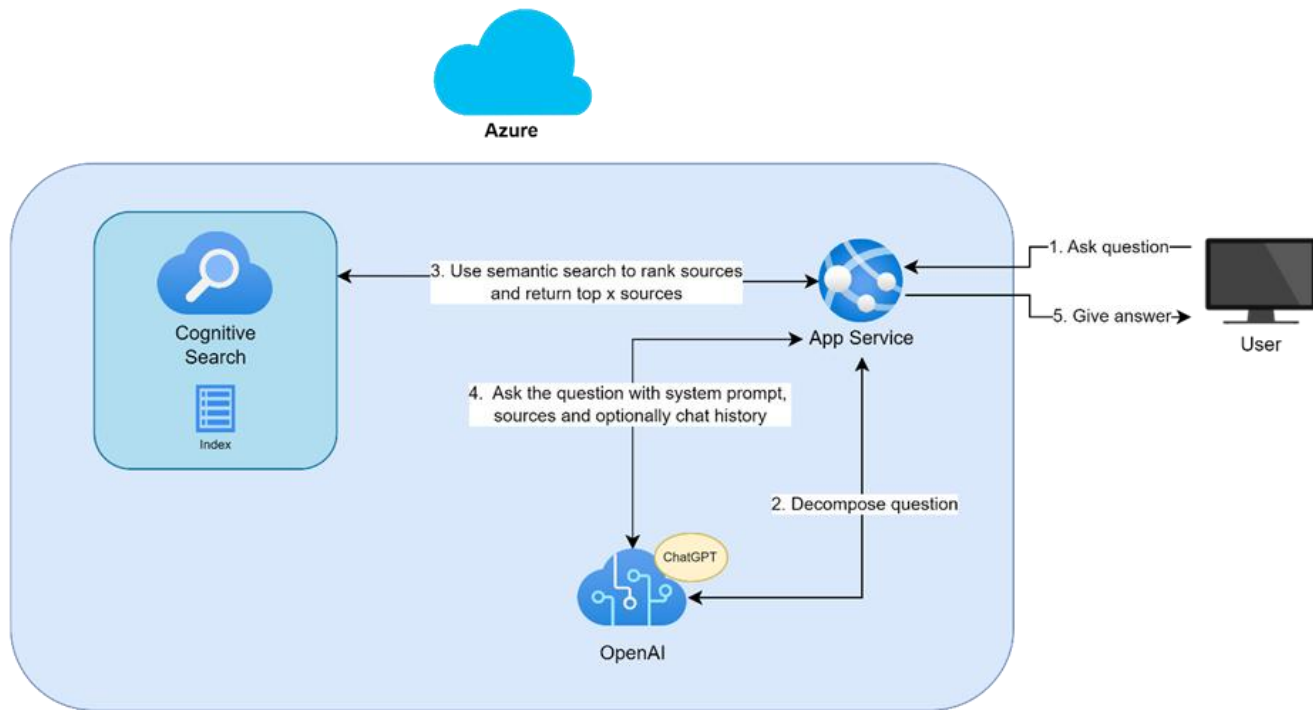
With RAG, we can find a (semantic similar) passage in your documents to answer a question with evidence from the text...



Azure AI Foundry building blocks are orchestrated to load documents into a semantic search index



An Azure AI Foundry large language model is used to answer the user's question based on the retrieved chunks from the semantic search index



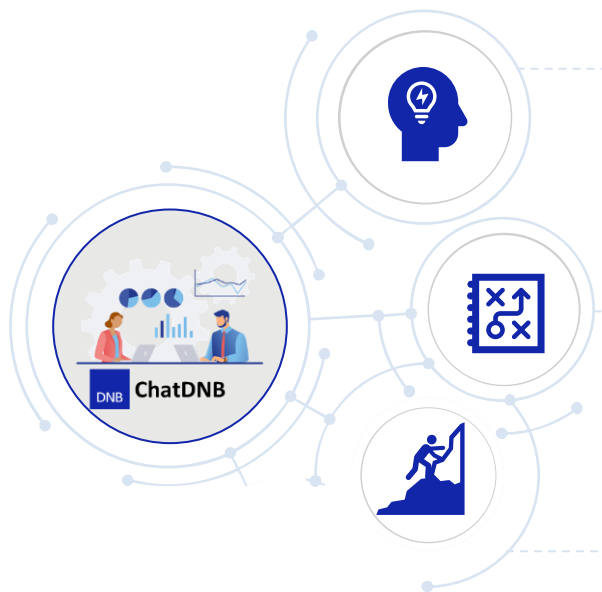
The quality of the answers can be improved by adjusting various parameters of the RAG pipeline

Pipeline part	Definition	How can it be adjusted?
 Data cleaning	The input text should be clean and not contain too many odd signs/figures/ tables. E.g., excel files	Select only relevant documents and clean the documents that contain too many icons after the document conversion step
 Chunking	Splitting texts into sentences, paragraphs or pages that can be vectorized (create an embedding)	Adjust the strategy to provide necessary parts of the text to answer your question. E.g., short questions vs. summarizing entire documents
 Text embeddings creation	Vectorization of DNB texts that are saved in the search index	Create vectors with accurate representation of DNB or domain specific jargon (e.g., "large exposures" vs. "big loans". Fine-tune a model or select a more capable one
 Embedding of the question	Vectorization of the question to make a comparison with the chunk vectors and retrieve the most relevant chunks	Educate users on how they ask questions or change the model to create different vectors of the question
 Document metadata	Data about the document that you are trying to search through	Specify the search area to look for input to the LLM on which the answer should be based. Saving the documents with correct labels, tags, or folders is essential.
 Retrieval parameters	Values to determine whether to retrieve a chunk and how many	Set values, given the use case preferences, to retrieve the correct number of chunks and decide on a similarity score for when an answer is deemed relevant
 Use of agents (LLM¹)	A LLM that performs various tasks within the pipeline	Use one or multiple agents for more complex tasks e.g., to check the relevancy of the chunks or to check if the evidence has really been used in the answer
 LLM	The LLM ¹ used to create an answer for your question	Select another LLM which can handle more context, is better in a specific domain or type of question, which is available given the cloud strategy and costs
 Prompt Engineering	The context and instructions that are given to the model to generate an answer	Create personas and a set of instructions to better answer the questions given the context



ChatDNB - Demo

Implementing this new technology proved to be a bumpy ride but we've gained many valuable insights throughout the process...



The vital role of an innovation officer to assist with:

- Identifying new opportunities
- Promoting adoption of the tool
- Keeping (core) users engaged and aligned in the development of new features

Given today's knowledge, we might have approached things differently:

- First looks matter: changing sentiment due to a disappointing first trial
- Direct clarity on the positioning of ChatDnB with respect to e.g., Copilot
- Identify needs: flexibility vs bulk
- Integrate in existing tools

We are still at the beginning of our journey of adopting GenAI:

- We need to continuously update workflows to AI workflows
- "AI will fix everything"
- We can expect more sophisticated applications that will change the way supervisors interact with all their data