



SME AI PROGRAM 2025 REPORT

From Documents to Answers

Enabling a Retrieval-Augmented Generation
(RAG) for Zurich SMEs.



Canton of Zurich
Department for Economic Affairs
Office for Economy



Swiss Data Science Center



Table of Contents

Table of Contents	2
Project Information	3
Executive Summary	4
Business Challenges	5
Knowledge Trapped in Documents.....	5
Sensitive Data and Sovereignty Requirements	6
Key Outcomes.....	6
Approach	7
Why RAG?	7
Why AI?	7
Why Collaborative Development?	7
The Shared Prototype	8
Results	8
Operational Insights: Companies Already Moving	8
Business Value	9
Technical Findings	10
What Other Zurich SMEs Can Learn	10
The Basis	10
Key Lessons	11
Technical Appendix	12
A. Use Case and Test Data Design	12
B. Development Structure: Baseline and Four Feature Tracks	13
C. Technology Stack	14
D. Core Design Principles.....	14
Open-Source Prototype	15
Related Information & Links	16



Project Information

Contributors

- Paulina Körner (Data Scientist), Ivan-Daniel Sievering (Senior Data Scientist), Thibaut Loiseau (Machine Learning Engineer) and Anna Fournier (Principal Data Scientist), Swiss Data Science Center (SDSC). More: datascience.ch
- Janine Plüss (CEO), Cibra. More: www.cibra.ch
- Marco Meola (CEO), DATABIOMIX. More: www.databiomix.com
- Adelene Lai (AI Integration Lead), Econetta AG. More: www.econetta.com
- Gosia Kubat (Manager and Co-owner) and Alma Thalmann (Manager and Co-owner), LanguageMasters GmbH. More: www.languagemasters.ch
- Roger Leimer (Head of IT), Microdul AG. More: www.microdul.com
- Wolfgang Loerli (CEO), SOORT AG. More: soort.com
- Detty Berta (Chairwoman of the Board), StratoServ Sciences AG. More: www.stratoserv.ch
- Marco Induti (Deputy CTO), Swiss Safety Center AG. More: www.safetycenter.ch
- Patrik Vonlanthen (Managing Director), Vonlanthen INSIGHT. More: www.vonlanthen.tv
- Markus Müller (Co-Head of the Division of Business and Economic Development) and Raphael von Thiessen (Program Lead AI Hub), Canton of Zurich. More: zh.ch/location
- Chantal Stäubli (CEO) and Anna Zakharova (Head Innovation Consulting), Ahead Zurich. More: ahead-zh.ch/en

Timeline

- September 2025 - March 2026: Roundtable collaboration towards AI prototype development
- April - May 2026: Transitioning AI prototypes to pilot projects

Executive Summary

Nine companies across manufacturing, consulting, health tech, environmental services, safety inspection, and commerce joined the Retrieval-Augmented Generation (RAG) track of the SME AI Program in the Canton of Zurich. Their sectors were different; their challenge was the same: valuable knowledge was already inside their organisations, within documents, manuals, reports, and assessments, but too hard to find, too slow to retrieve, and too scattered across sources to synthesize into a reliable answer.

Over two weeks of collaborative development with the Swiss Data Science Center (SDSC), the cohort built a shared, modular RAG prototype, grounded in a realistic compliance scenario and extensible to each company's own use case.

Prototype adoption: several companies transitioned the prototype to live deployment within weeks of the program concluding.



“The most valuable learning was confirming that AI is nothing without the right data. We started with a specific use case, but soon realized that the real challenge was not the technology itself. Questions around data availability, accessibility, ownership, and confidentiality quickly emerged. We also learned that while approaches such as RAG are already well established, successfully implementing them inside an organization requires much more than connecting an AI model to a database.”

— Marco Induti, Deputy CTO, Swiss Safety Center AG

The central lesson across the cohort: data quality, not the AI model, is the first bottleneck every company encountered. Understanding the full pipeline, rather than treating it as a black box, is what allows a company to keep improving it after the program ends. Collaboration across sectors and maturity levels produced more grounded outcomes in two weeks than any single company could have reached alone in the same timeframe, and the companies that defined their business metrics before deployment, not after, are the ones now able to show it.

The developed RAG pipeline is open-sourced and can be found here: <https://github.com/SwissDataScienceCenter/sme-kt-zh-collaboration-rag>

Business Challenges

The nine participating companies ranged from those just beginning their AI journey to those already running machine learning in production, but they arrived with the same structural problem, not a technology request. In each case, knowledge that already existed inside the organisation was either too slow to access, too dispersed to synthesise, or too sensitive to expose to external services.

- **Knowledge trapped in documents:** years of accumulated reports, manuals, and assessments existed electronically but required expert time to find and interpret.
- **Data sovereignty:** several companies handle confidential or regulated data, ruling out standard cloud AI tools.
- **Reliable retrieval:** the answers generated from the documents should be traceable to their source, making verification possible and reducing the risk of acting on incomplete or misinterpreted information.

The specific domains differed; the structural pattern did not. Three companies illustrate how this played out in practice.

Knowledge Trapped in Documents

Microdul AG specializes in miniaturized electronics for medical devices and industrial applications, operating in a heavily regulated environment. Engineers routinely consult a large base of design documents, validation records, production instructions, and component specifications. Finding the right document reliably, and detecting when information in one document contradicted another, was both time-consuming and a compliance risk.

“For ensuring our seamless production, our engineers perform quite a large volume of sophisticated and often repeated design and production instruction work. A large base of knowledge-heavy documents is being created during this process.”

— Daniel Thommen, Head of Module Engineering, Microdul AG

And adding to that point:

“[...] The first challenge is to find and retrieve the correct information from this large information base. Second, automated analysis and inconsistency detection would help reduce repeated work and enforce high quality levels.”

— Roger Leimer, Head of IT, Microdul AG

Econetta AG, an independent environmental consulting firm, faced a different version of the same problem. Project teams needed access to regulatory knowledge, project history, and technical analysis spread across siloed, mostly unstructured files. Accessing that knowledge took time that should have gone to client work and the ability to answer certain questions depended too heavily on individual team members who held the knowledge informally.

Swiss Safety Center AG, an accredited inspection and certification body, collects large volumes of inspection data through its daily work. Beyond routine retrieval, the company saw potential value in detecting patterns and correlations across cases that manual review could miss, risk signals that could inform safer outcomes, but had no infrastructure to surface them systematically.



Sensitive Data and Sovereignty Requirements

Several companies handled data that could not leave their own infrastructure. **Vonlanthen INSIGHT**, a boutique Swiss consultancy, works with confidential client documents. The requirement was a privacy-preserving RAG system that runs entirely on local hardware, with no data sent to external providers, but still delivers answers reliable enough for professional client engagements.

“[...] I needed a way to extract structured insights from large, heterogeneous document sets without sending proprietary data to cloud providers. The challenge was building a privacy-preserving RAG system that could run entirely on local infrastructure, handle real-world document complexity, and deliver answers reliable enough to use in professional client engagements.”

— Patrik Vonlanthen, Vonlanthen INSIGHT

DATABIOMIX, a health tech company developing AI infrastructure for gut microbiome analysis, faced a scale challenge: thousands of heterogeneous scientific publications needed to be synthesised into evidence-based dietary recommendations with citations that clinicians and nutritionists could verify and trust in their practice.

The other four companies in the cohort, **Cisra**, **LanguageMasters GmbH**, **SOORT AG**, and **StratoServ Sciences AG**, faced variations of the same knowledge-access and trust problems. Their specific use cases are not detailed here.

Key Outcomes

9 companies built a shared RAG prototype together, fully open-source and freely available to any Swiss SME.	5 companies actively deploying or extending the prototype within weeks of the program.	2 production client applications launched by Vonlanthen INSIGHT based on the prototype.
--	---	--

Approach

Why RAG?

Retrieval-Augmented Generation (RAG) is built for one specific purpose: making existing knowledge accessible through natural language, with cited sources. Unlike a general AI assistant that draws on pre-trained knowledge, a RAG system retrieves from your own documents and grounds every answer in what is actually there.

This made RAG the right fit for these companies for three reasons:

- The knowledge already existed inside the organisation. The goal was access, not generation.
- Citations matter. In manufacturing compliance, safety inspection, consulting, and healthcare, an answer without a verifiable source is not actionable and carries professional risk.
- An honest “I don’t know” is safer than a confident wrong answer. A system that fabricates answers is more dangerous than one that acknowledges the limits of its evidence. RAG architecture can be specifically designed to cite the underlying sources and acknowledge when the evidence is missing.

Why AI?

The volume and heterogeneity of documents across these companies had already exceeded what is practical to manage manually. Traditional keyword search returns documents, not answers. AI-powered retrieval returns grounded, cited responses to natural language questions, reducing the gap between a question and a decision. For companies with thousands of documents, this is qualitatively different from anything available before.



“Seeing how SMEs from completely different industries framed their AI challenges was extremely helpful. The underlying patterns – data quality, change management and most importantly: trust – are universal, really independent of the sector.”

— Patrik Vonlanthen, Vonlanthen INSIGHT

Why Collaborative Development?

Nine companies built together on a shared codebase rather than each developing in isolation. Using modern development tools we could use each company’s strength to efficiently create a sound solution and gain efficiency throughout the program.

Companies at earlier stages of AI adoption saw first-hand what production-grade AI actually requires: not just a working prototype, but data governance, evaluation infrastructure, and a plan for maintenance. More experienced companies were pushed to articulate design choices they had previously taken for granted.

The shared prototype also encountered failure modes that no single company would have surfaced alone. A system designed for regulatory compliance documents behaved very differently when tested against scientific literature or supplier self-declarations and those differences produced design improvements that benefited everyone.

The Shared Prototype

The prototype was grounded in a concrete sustainability compliance scenario for a fictional packaging company which answered questions about supplier claims, environmental certifications, and procurement policy. This gave every design decision a testable, realistic context. The architecture covered the full RAG pipeline: document ingestion and indexing, retrieval, answer generation, and a working chat interface with source citations. The fictional context with real documents enabled seamless collaboration and a shared vision among participants, while still enabling each company to pursue their own angle.

Development was structured around a baseline pipeline and four progressive feature tracks: evaluation, structured outputs, advanced retrieval, and agentic reasoning, allowing nine companies with different priorities to contribute to and benefit from the same architecture simultaneously. All code is published open-source and publicly maintained by the SDSC on [GitHub](#), and free for any Swiss SME to use and adapt.

Results

Operational Insights: Companies Already Moving

The clearest signal from the program was the speed at which companies moved from prototype to application. Several began within weeks of the collaborative sessions, without waiting for additional support.

Vonlanthen INSIGHT progressed far. Within months, the SDSC prototype had grown into TRAG, an open-source RAG framework for systematic parameter testing and production deployment. Two client projects now run on it: one for conversational access to personal document archives, one for on-premise document indexing for public sector and SME clients where local processing ensures full data sovereignty.



“The SDSC-provided prototype I worked with during the program has since grown into TRAG on Github: github.com/voninsight/TRAG. A RAG framework designed for systematic parameter testing and production deployment. Two of my actual projects build on TRAG: ONE, an experimental personal knowledge cockpit for conversational access to one’s own document universe; and Archivio, an on-premise indexing and querying system for government bodies and SMBs – local models for full data sovereignty, or cloud APIs where permitted.”

— Patrik Vonlanthen, Vonlanthen INSIGHT

Microdul AG forked the program repository immediately after the sessions ended. Two participants from the company extended it with production-oriented features - support for additional model backends and basic user role management. Deployment alongside an internal language model is now underway.

“With the app ready, the actual implementation is now underway. Deploying an internal LLM that is used in combination with the RAG system unlocks more potential for innovation [...] Using standardized APIs, the LLM can be used either standalone using a chat interface, or be wired into other digital processes.”

— Roger Leimer, Head of IT, Microdul AG

DATABIOMIX is integrating the RAG layer directly into its web application, enabling nutritionists to translate gut microbiome profiles into evidence-based, citable dietary recommendations.

Swiss Safety Center AG recognised during the program that successful deployment requires dedicated internal ownership. The company hired a data scientist and is now designing a full retrieval pipeline for integration with its inspection knowledge base.



“The program helped us understand what is technically possible and what is required behind the scenes to make it work. The next step is to bring these learnings into our own environment. We are currently facing the typical challenges of integration, governance, and limited AI resources, but our goal is to deploy a first working agent connected to our internal knowledge base. From there, we want to implement the full retrieval pipeline we have designed and make it stable, measurable, and maintainable before validating it with real users and use cases.”

— Marco Induti, Deputy CTO, Swiss Safety Center AG

Econetta AG is working on multiple RAG use cases, combining internal memo and regulatory advice with European legislations. They are building pipelines to ingest external public data and tracking regulatory developments.

“This program gave us a tangible starting point to incorporate AI as a solution, both intellectually and practically. Intellectually, the demystification sessions gave us a good understanding of RAG at an applied level; practically, we could start prototyping immediately using the MIT-licensed RAG codebase SDSC shared. Compared to other open RAG solutions available that are often ‘black box’, SDSC’s solution is open, modular and comes with educational materials to help us troubleshoot failure modes.”

— Adelene Lai, Econetta AG

Business Value

Across the cohort, the program delivered business value in three concrete forms:

- **Capability acceleration:** in as little as two weeks of prototyping, companies gained a validated prototype and the foundational understanding to extend it, work that would otherwise have taken months of trial and error.
- **Zero-cost infrastructure:** the open-source prototype can be deployed by any Swiss SME at no licensing cost, directly lowering the financial barrier to AI adoption.
- **Faster internal adoption:** having a working prototype, and understanding how it was built, gave companies the evidence and vocabulary to secure internal buy-in for continued AI investment.

The companies that completed the program leave with everything required to generate real business impact: a working prototype they built and understand from the inside, and the organisational moves already made, like new hires, new roles, and active deployments. The question is no longer whether AI can work in their context, but how quickly the systems now in place translate into time recovered, decisions improved, and knowledge that was previously locked in documents made genuinely accessible.

“The real output of this program is not the prototype alone. It’s what companies built alongside it: clearer AI roadmaps, budget allocated to AI initiatives, and people hired into AI roles. Several companies have already confirmed further developments, which is why the momentum has not stopped now that the first round of the SME AI program has closed.”

— Anna Fournier, Swiss Data Science Center



Technical Findings

Three design insights proved consistently important across all use cases. Full technical detail is in the Technical Appendix.

- Retrieval quality is not controlled by any single decision. How documents are split, how they are indexed, and how queries are formulated each contribute to final answer quality. No single step dominates, instead all steps compound.
- Without a way to measure quality, it is difficult to know which changes help and which do not. Adding an evaluation step early, before layering on more features, makes each improvement legible and prevents problems from compounding silently.
- The workshop tested semantic search, keyword search, and hybrid combinations. Hybrid retrieval is designed to handle both vocabulary mismatch and exact-term queries, making it more consistent across different query types, though results vary by use case.

What Other Zurich SMEs Can Learn

The Basis

An AI system is only as useful as the information it draws on. Documents that are scanned rather than digital, inconsistently structured, out of date, or stored in incompatible formats will require significantly more preparation work before any AI system can use them reliably. The quality of your data determines how much effort stands between you and a working system. This is not a reason to delay. It is a reason to include a data quality assessment as the explicit first step of any AI project, even before any vendor selection or technical work begins.

Organizational readiness matters just as much: a working prototype turns into a liability the moment nobody owns the documents, the access rights, or the approval of what the system produces. Swiss Safety Center AG's decision to hire a dedicated data scientist before deployment reflects exactly this: a technical solution needs a technical owner.

The initial deployment is a starting point. Retrieval quality improves as documents are cleaned, prompts are refined, and edge cases are addressed. Plan for iteration, not installation. Companies that treated AI as a developing capability progressed consistently further than those seeking a one-off solution. The field itself reinforces this: models, tooling, and retrieval techniques are advancing quickly, and an organisation that stays engaged is better positioned to benefit from those improvements than one that considers the project closed.

Is RAG relevant for your business? Start with these questions.

About your knowledge and documents:

- ⇒ Where do your employees spend the most time searching for information?
- ⇒ Which decisions require synthesising multiple documents before acting?
- ⇒ Does critical knowledge depend on specific individuals, what happens when they are unavailable or leave?

About privacy and data:

- ⇒ What data can you share with a cloud AI provider, and what must remain on-premise?
- ⇒ Are there regulatory or contractual constraints on how your documents may be processed?
- ⇒ Who owns the documents and has authority to grant AI access to them?

About readiness:

- ⇒ How consistent, current, and well-structured are your internal documents?
- ⇒ Who in your organisation would own and maintain an AI knowledge system?
- ⇒ Can you define what a good answer looks like before you deploy?

Key Lessons

- **Start with the business problem, not the technology.** Every company that made meaningful progress began by naming a concrete operational frustration, not a technology aspiration. "Where are we losing time to document search?" and "where does missing information cost us most?" are more productive questions than "how can we use AI?" The technology choice follows from the answer, not the other way around.
- **Define what success looks like before you start.** Naming the problem is not the same as knowing whether you've solved it. The companies that progressed to pilot most quickly, had a clear sense of which documents the system needed to work with. Pairing that clarity with a measurable baseline makes the case for continued investment much stronger. That baseline can be simple: how long does it currently take an employee to answer a typical query, or how often do errors occur today? Captured before deployment, a number like that gives you something concrete to compare against later. Several companies said they would build this measurement approach even earlier in a future project.
- **Data preparation is where the real effort goes, and where it pays off.** Across the cohort, the work of cleaning, structuring, and standardising documents took more time than anticipated, and it directly shaped what the system could do. The payoff: a well-structured document base with a straightforward retrieval approach consistently outperformed a more sophisticated model working on poorly prepared data.
- **RAG-specific: the pipeline is the product.** Building a RAG system means making a sequence of interconnected decisions, how documents are ingested and chunked, which method converts them to searchable form, how queries are phrased, how results are ranked, and how the final answer is produced. Each decision affects the quality of the next. Companies that understood the full pipeline could diagnose failures and improve methodically; those that treated the system as a black box could not.
- **In professional contexts, a grounded 'I don't know' is more valuable than a confident wrong answer.** For companies in regulated industries, compliance work, or client-facing roles, an AI that declines to answer when evidence is insufficient is a feature, not a failure. The shared prototype was explicitly designed to distinguish verified facts from unverified claims, and to refuse to answer when the evidence did not support a conclusion. This does not eliminate hallucination risk entirely, but it reduces it significantly and makes the system's uncertainty visible rather than hidden. For high-stakes contexts, that transparency is as important as the answer itself.



"[...] A RAG is not one model but a full pipeline, and the real work lives in the steps between. Going from raw corpus to production means making and tuning countless decisions: chunking strategy, embedding model, retrieval and reranking, prompt design, and the parameters to set at each stage. Seeing how every choice compounds into final output quality was the biggest shift in how I approach building one."

— Marco Meola, Co-founder and CEO, DATABIOMIX

"What stands out from this cohort is that nine companies, from different sectors and different stages of AI maturity, collaboratively built a solution more comprehensive than anything a single SME would have created alone."

— Raphael von Thiessen, Canton of Zurich



Technical Appendix

This appendix is for readers with implementation interest. It summarises the design decisions and architecture of the shared prototype. The full codebase, including notebooks and evaluation datasets, and a detailed README covering installation, feature tracks, and repository structure, is available at: github.com/SwissDataScienceCenter/sme-kt-zh-collaboration-rag

A. Use Case and Test Data Design

The prototype was built around a sustainability compliance scenario for a fictionalised packaging company (PrimePack AG), designed specifically to test the system against the failure modes that naive chatbots exhibit. The document corpus (~40 documents) included: third-party verified Environmental Product Declarations (EPDs); supplier datasheets with self-declared sustainability claims; reference standards (GHG Protocol, CSRD guidelines, ISO standards); and deliberately flawed artificial documents containing missing data, unverified claims, and temporal conflicts.

The flawed documents were included to assess whether the system could distinguish evidence quality and communicate the limits of its knowledge — testing for hallucination (responses about out-of-scope products), overconfidence (handling of unverified claims), fabrication (behaviour when data is absent), and silent error (resolution of conflicting sources). A ground-truth evaluation dataset was prepared covering five categories: portfolio scope, claim verification, missing data, source conflict, and policy compliance.

B. Development Structure: Baseline and Four Feature Tracks

Baseline RAG Pipeline

The six core stages: ingest (load and parse raw documents into a common format), chunk (split documents into retrievable segments), embed (convert to vector representations), store (index in a vector database), retrieve (find relevant segments for a query), and generate (produce a grounded answer). Two document parsers were tested and compared; chunking strategies were evaluated for different document structures and content types.

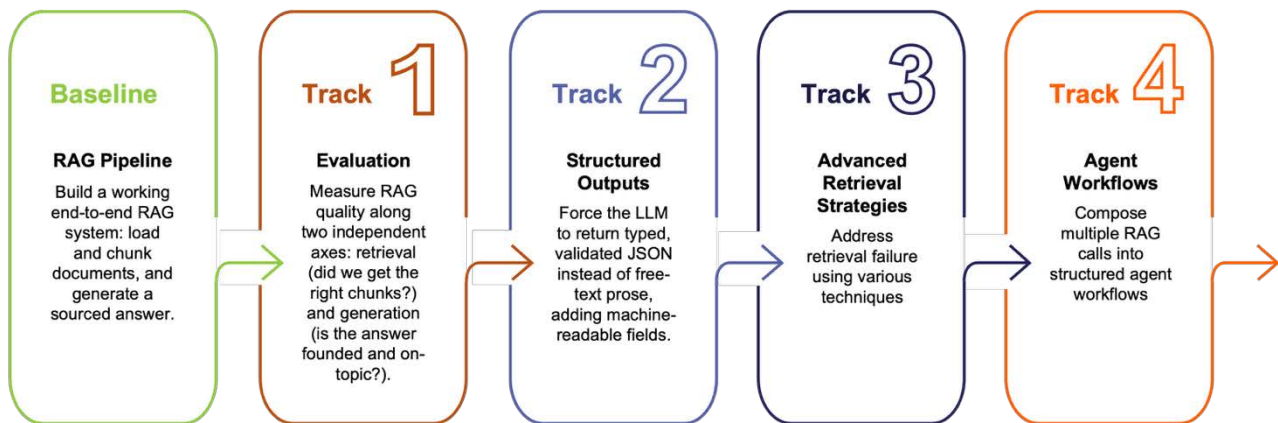


Figure 1: RAG prototype development structure – baseline pipeline and four feature tracks, SME AI Program in the Canton of Zurich (February-March 2026)

Track 1: Evaluation and Multi-Modal Retrieval

The RAGAS evaluation framework was integrated, providing four quality metrics: faithfulness (does the answer contradict its sources?), answer relevancy, context precision, and context recall. A synthetic question-answer dataset was generated from the corpus to enable evaluation at scale. Multi-modal retrieval was also added: images were embedded and stored in a separate vector collection, enabling retrieval of visual evidence such as product images, labels and certification stamps alongside text.

Track 2: Structured Outputs

JSON-schema-constrained outputs were introduced so that every response includes a structured answer field and an explicit list of source document identifiers. This makes outputs auditable and suitable for integration into reports, dashboards, or downstream workflows.

Track 3: Advanced Retrieval

Three retrieval strategies were implemented and compared: vector similarity search (semantic matching), BM25 keyword search (exact term matching), and hybrid retrieval combining both using Reciprocal Rank Fusion. Query intelligence techniques added include query expansion (generating multiple phrasings of the same question before retrieval) and HyDE (generating a hypothetical answer and using it as the retrieval query). A conversational context manager was added for multi-turn coherence.

Track 4: Agentic Workflows

The system was extended toward multi-step reasoning: the language model was given access to retrieval as an explicit tool call, with an agent loop that allows it to reason, call tools, observe results, and continue before producing a final answer. This enables the system to decompose a complex question into sub-queries, check multiple sources independently, and synthesise a grounded response. A RAG-as-subagent pattern was also demonstrated for modular integration

C. Technology Stack

The prototype was built entirely on established open-source components to ensure accessibility, reproducibility, and freedom from vendor lock-in:

- Embedding models: converting text (and images) to vector representations for retrieval
- Vector databases: indexing and storing document embeddings
- Large Language Models (LLMs): both locally hosted and cloud-hosted models were tested and compared
- Renku: SDSC's collaborative computing platform for shared notebooks and demo hosting
- GitHub: version control, collaboration, and public distribution of code and documentation

The codebase is organized as an editable Python library (conversational-toolkit) alongside an application layer that adapts it to a specific use case. The code is open-sourced and publicly maintained at:

github.com/SwissDataScienceCenter/sme-kt-zh-collaboration-rag

D. Core Design Principles

Three principles drove every architectural decision and apply regardless of domain or technology stack:

- **Separate evidence levels.** Distinguish verified facts (e.g. third-party certified EPDs) from unverified claims (e.g. supplier self-declarations). In regulated or compliance contexts, this distinction is not optional.
- **Build evaluation infrastructure early.** Measuring quality before adding features enables incremental improvement and early detection of failure modes. Building evaluation after the fact is significantly harder.
- **Structure outputs for auditability.** Free-text answers are difficult to verify; structured outputs with explicit source citations enable downstream validation, integration into workflows, and the confidence that comes from traceability.

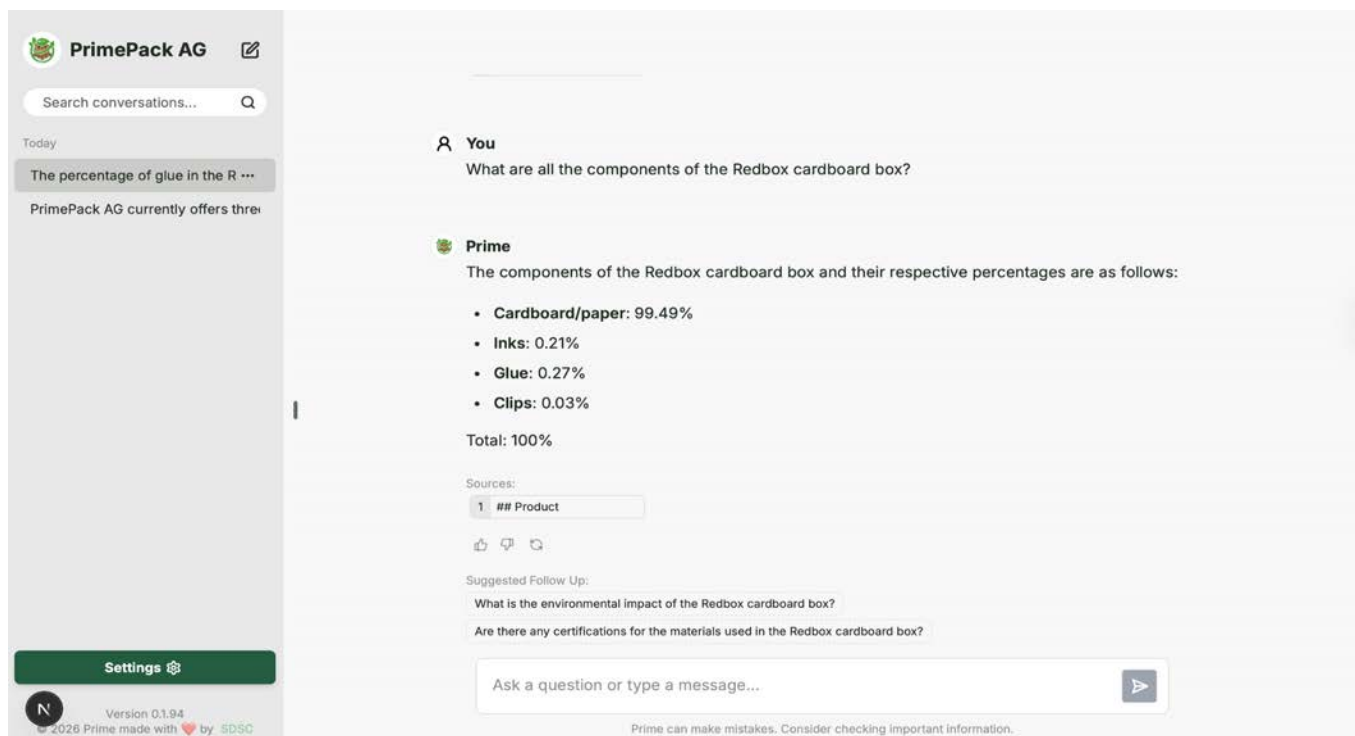
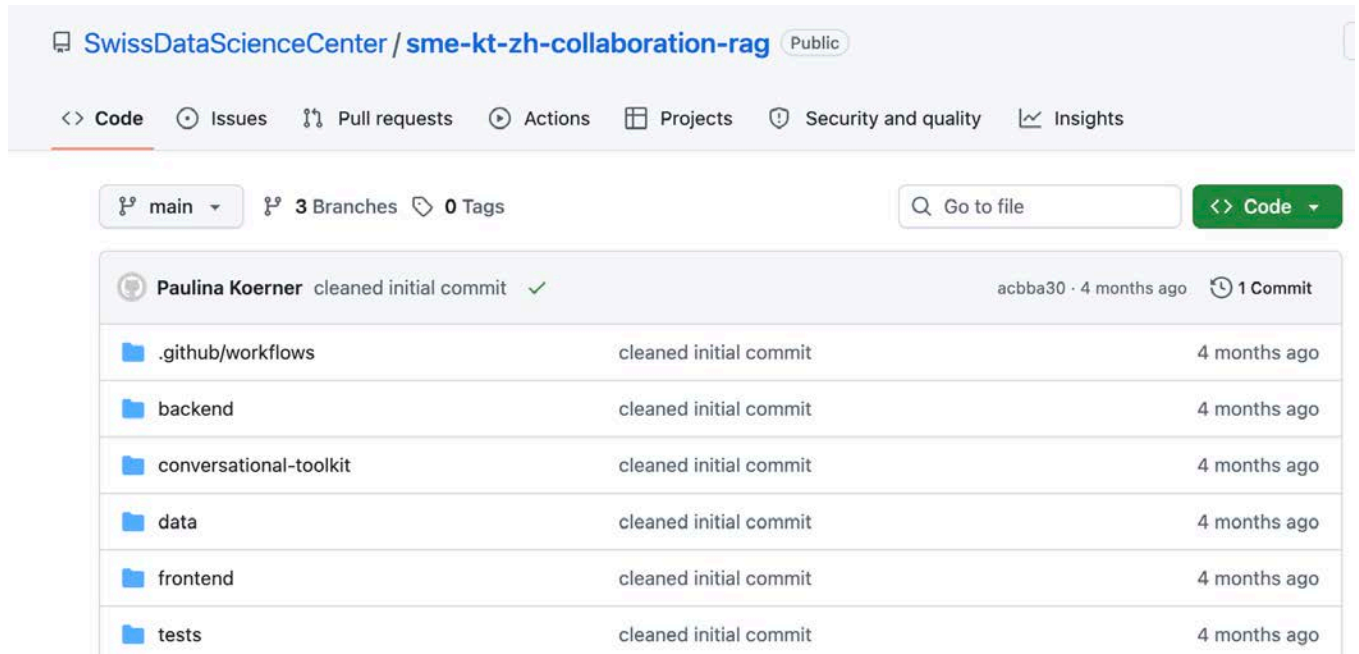


Figure 2. RAG prototype frontend (PrimePack AG use case), SME AI Program in the Canton of Zurich (February–March 2026).

Open-Source Prototype

The RAG prototype is available open-source and maintained by the SDSC: <https://github.com/SwissDataScienceCenter/sme-kt-zh-collaboration-rag>



SwissDataScienceCenter / sme-kt-zh-collaboration-rag Public

<> Code Issues Pull requests Actions Projects Security and quality Insights

main 3 Branches 0 Tags Go to file Code

Paulina Koerner cleaned initial commit ✓ acbba30 · 4 months ago 1 Commit

.github/workflows	cleaned initial commit	4 months ago
backend	cleaned initial commit	4 months ago
conversational-toolkit	cleaned initial commit	4 months ago
data	cleaned initial commit	4 months ago
frontend	cleaned initial commit	4 months ago
tests	cleaned initial commit	4 months ago

Related Information & Links

SME RAG Systems – Watch Video Case Study:

https://www.youtube.com/watch?v=wtDOt9v_ngs



Related Links:

[Canton Zurich SME Program – SDSC](#)

[Open-source RAG prototype on GitHub](#)

zh.ch/location

innovation.zuerich

ahead-zh.ch

About the Division of Business & Economic Development at the Office of Economic Affairs, Canton of Zurich

We promote innovation and support businesses – from founding to networking within Zurich's key industries. Through our platform 'Innovation.Zurich', we provide guidance and information about the region as a hub for innovation. Together with the consortium Ahead, we assist local companies in pursuing their innovation projects. We also serve as a point of contact for businesses interested in establishing operations in Zurich. More information: zh.ch/location | innovation.zuerich | ahead-zh.ch

About the Swiss Data Science Center (SDSC)

The Swiss Data Science Center (SDSC) is a national research infrastructure in data science, machine learning and artificial intelligence (AI), founded by EPFL and ETH Zurich. Its mission – to enable data-driven science and innovation for societal impact – drives its initiatives in research projects, knowledge and technology transfer, and education. With a large multidisciplinary team of professionals in Lausanne, Zurich and at PSI in Villigen, the SDSC provides expertise and services to various domains, such as health and biomedical, energy and sustainability, climate and environment, digital society, and large-scale scientific infrastructures. The SDSC also contributes to initial and executive education programs at EPFL and ETH Zurich. More information: www.datascience.ch.

Contact:

For questions about the program, please contact:

Dr. Anna Fournier

Principal Data Scientist

Swiss Data Science Center

anna.fournier@datascience.ch

We thank the Zurich SMEs of the 2025 AI Innovation Program Cohort:



cisra.ch



databiomix.com



econetta.com



gonina.com



languagemasters.ch



microdul.com



papillon.ch



soort.com



stratoserv.ch



safetycenter.ch



vonlanthen.tv

The SME AI Innovation Program in the Canton of Zurich is a joint initiative of:



**Canton of Zurich
Department for Economic Affairs
Office for Economy**

zh.ch/location



Swiss Data Science Center

datascience.ch