

Student GrantMatch

Team & Group: G4, GardenStateAI

Atishay Jain, Dinesh Chandramohan | atishay23@gmail.com, dineshcm01@gmail.com

Repo: <https://github.com/atishaydottech/Capstone-Project> | Approach: Agentic RAG with Pydantic AI

Contents

1. Problem Statement & Goals 2. Architecture 3. Tool Stack 4. Evaluation Criteria

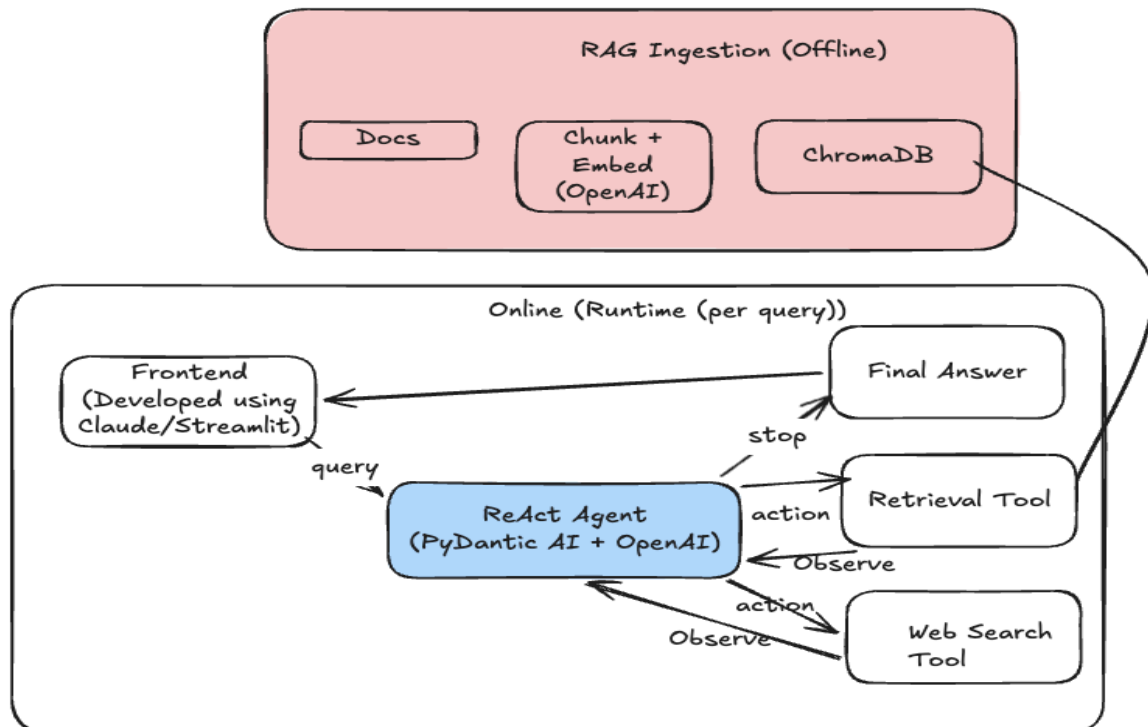
1. Problem Statement & Goals

Students often don't know which scholarships or grants they actually qualify for. Rules are scattered across long documents, and keyword search misses matches or gives wrong ones. This tool takes a student's basic info and tells them which programs they qualify for, citing the exact rule that says so.

Goals

- Check a student's info (income, state, GPA, major, etc.) against real program rules.
- Show the exact rule behind each yes or no, not just an answer.
- Flag partial matches instead of forcing a plain yes or no.
- Search documents only when needed, not on every query.
- Test against labeled profiles and measure real accuracy.
- Report eligibility only — not odds of winning.

2. Architecture



3. Tool Stack

- Pydantic AI: agent loop and tool-calling (Thought / Action / Observation), plus structured, validated outputs.
- OpenAI API: LLM for reasoning + embeddings (text-embedding-3-small).
- ChromaDB: local vector store for retrieved chunks.
- Retriever tool: wraps the vector store as a callable agent tool.
- Web search tool (Tavily / SerpAPI): second tool, gives the agent a real choice.
- Tracing: logs tool calls and latency for evaluation.
- Streamlit or Claude-made UI: simple front end for the demo.

Framework reasoning: Pydantic AI provides the ReAct-style loop needed here without a second framework on top, and its structured output format (a typed schema for eligibility + cited clause) directly supports precision and citation-accuracy scoring — avoiding fragile parsing of free-form text.

4. Evaluation Criteria

Quantitative

- Eligibility accuracy vs. 25–30 hand-labeled student profiles.
- Citation accuracy: correct rule quoted, not just a correct answer.
- Hallucination rate: cited details not actually in the source.
- Tool-call count and response latency per query.
- Consistency across repeated identical queries.
- Targets: >80% eligibility accuracy, >70% citation accuracy, <10% hallucination rate.

Qualitative

- Rate cited rules as correct, close, or unrelated on a sample.
- Review worst-performing cases and tag the failure cause.
- Check answers are understandable without reading the source document.
- Check tone is neutral and appropriately cautious on uncertain cases.
- Manually review borderline/edge-case profiles for nuance.