



STARS-GPT

(AI Tutoring Chat Agent)

Team Members: Vincent Carrancho, Robin Obregon, Stephanie Torres, Nicholas Belschner, Garvee Valcourt

Product Owner: Dr. Patricia McDermott-Wells



Problem Statement

- Since LLMs have been made publicly available, students have started using them for help with their assignments.
- LLMs are powerful tools for education. However,
 - LLMs are inconsistent sometimes and can lead to wrong results/solutions.
 - LLMs are often used by students to do their work for them.
 - LLMs are trained for general purposes and not something specific as schoolwork.
- This has been a rising concern for FIU's STARS Tutoring group.





Our Solution

- A custom application for the STARS Tutoring Group that address the problems of LLMs by:
 - Tailoring the responses of the LLM to specific subject matter.
 - Letting the LLM lead the student/tutee into the right direction.
 - Fine-Tuning the responses from the LLM





Different Use Cases

- Generation of examples for user to work through.
- Assignment/Exam Preparation
 - Clarifying concepts.
 - Assistance with coding assignments and debugging.
 - Offering in further in-depth information about the specific subject matter.
- And much more!





The Application

FIU

Engineering
& Computing

FLORIDA INTERNATIONAL UNIVERSITY



Requirements

Requirements

- Had to use OpenAI's API.
 - The LLM cannot do the work for the user (in completing assignments, etc).
- Had to have a user-friendly interface.
- Had to be widely accessible.
- Must be embeddable in Microsoft Teams.

Technical Requirements

- Database to store user information/messages.
- Application must be accessible.
- Must be secure.





We decided to use these Technologies

Frontend

- React.js
 - Library for creating interactable user interfaces.
- ShadCN
 - Ready to use accessible components.
- Tailwind CSS

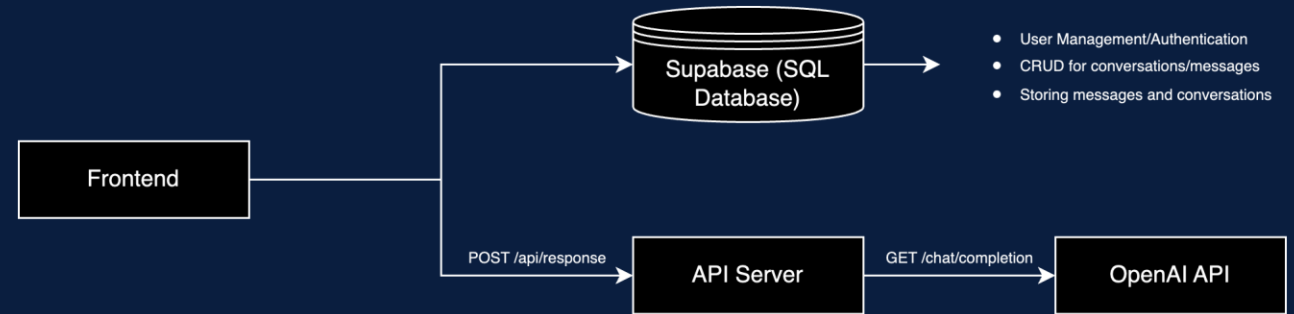
Backend

- FastAPI
 - Server to interact with OpenAI's model
- Supabase
 - Authentication Provider
 - Used to store user information, messages, conversations, etc.
- Render.com - Website Hosting
- AWS EC2 Instance for backend hosting (will change to Azure VMs in the future).

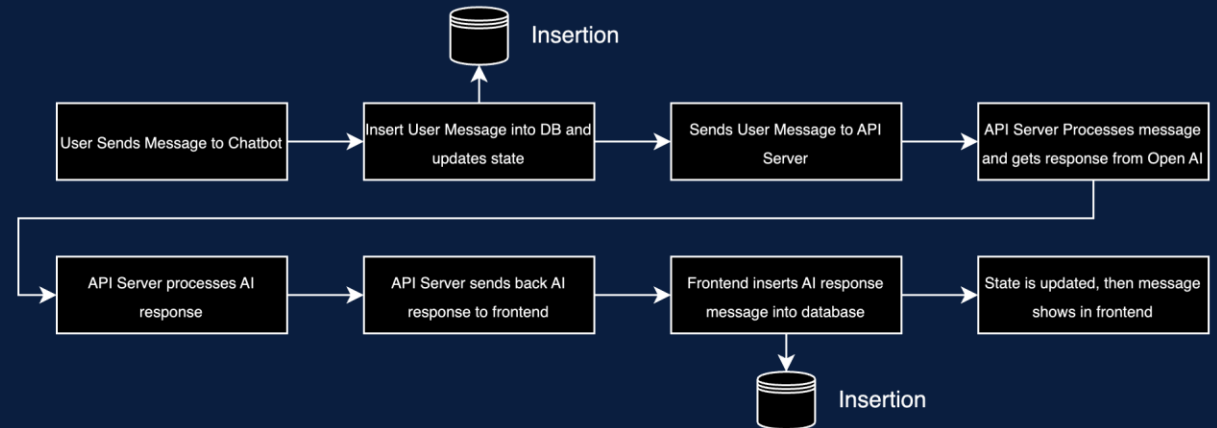


System Design

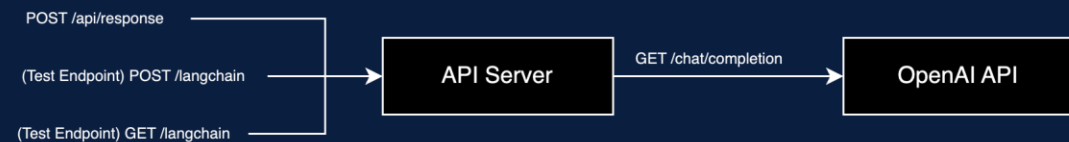
System Architecture:



AI Messaging System:

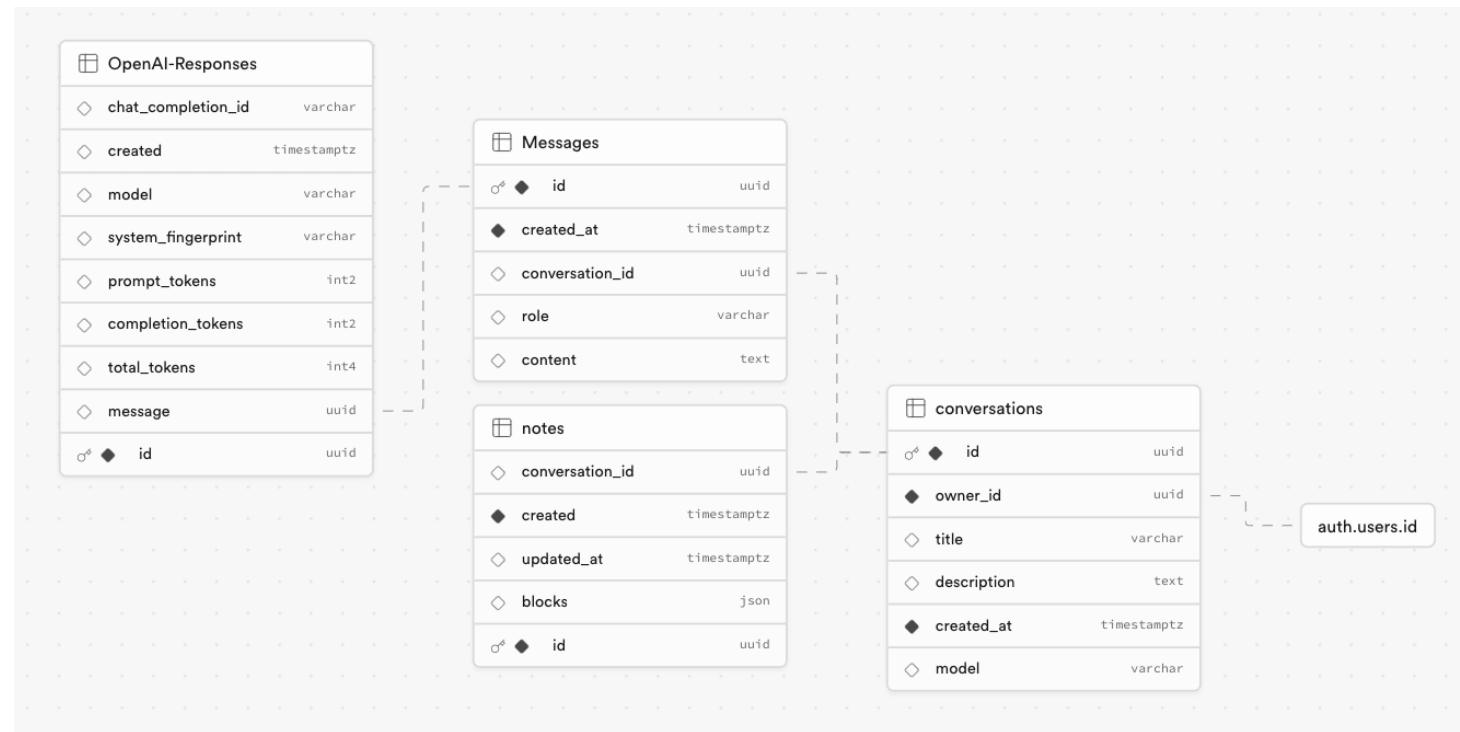


API Server Design

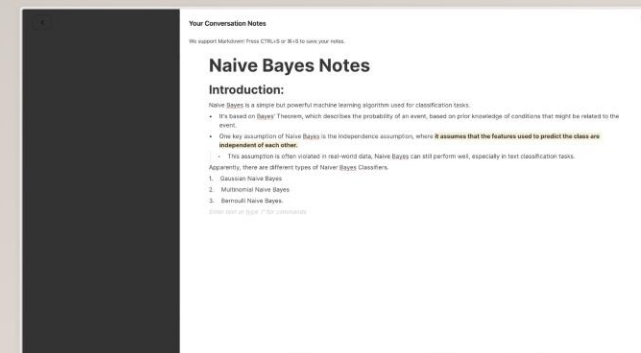
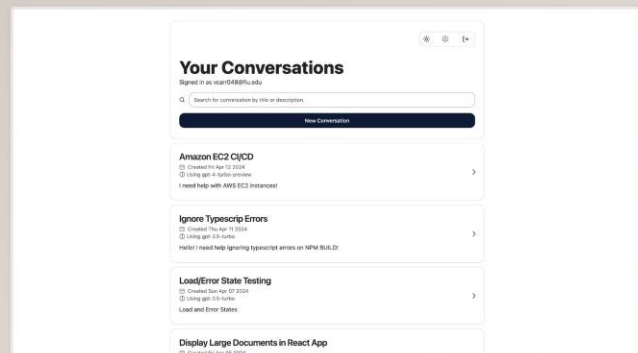
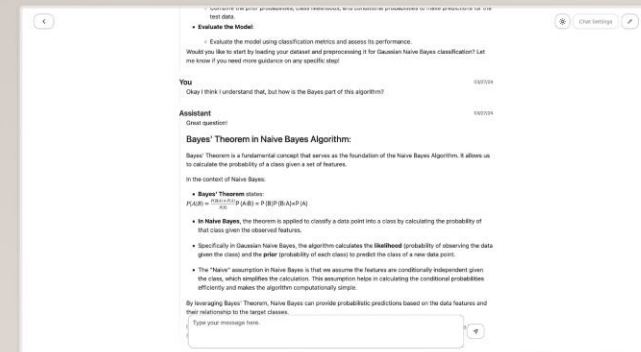
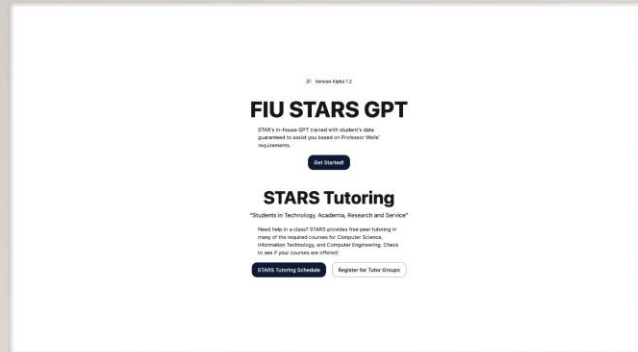


Database

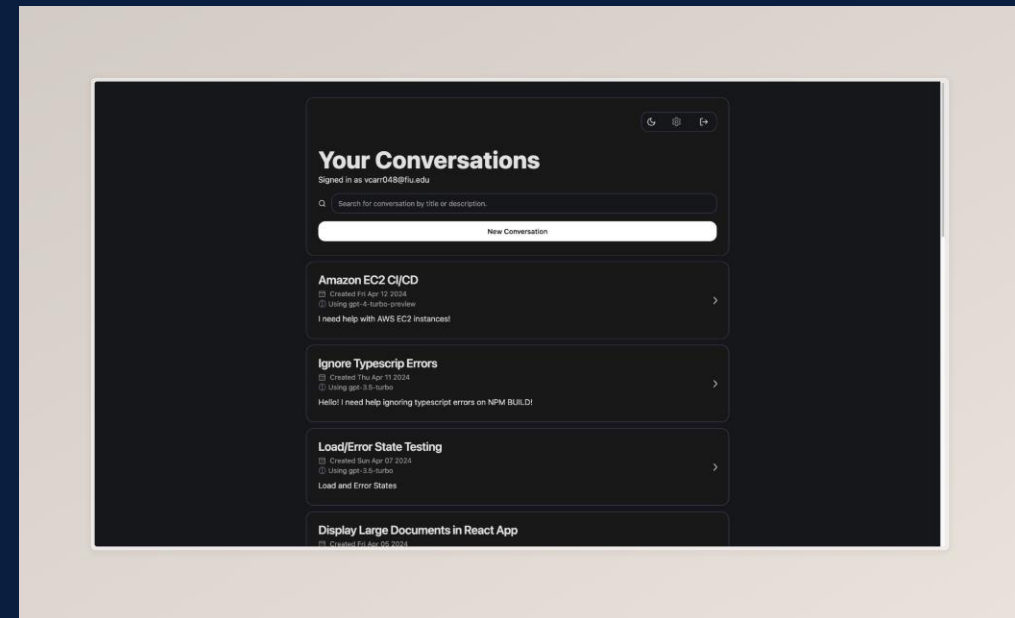
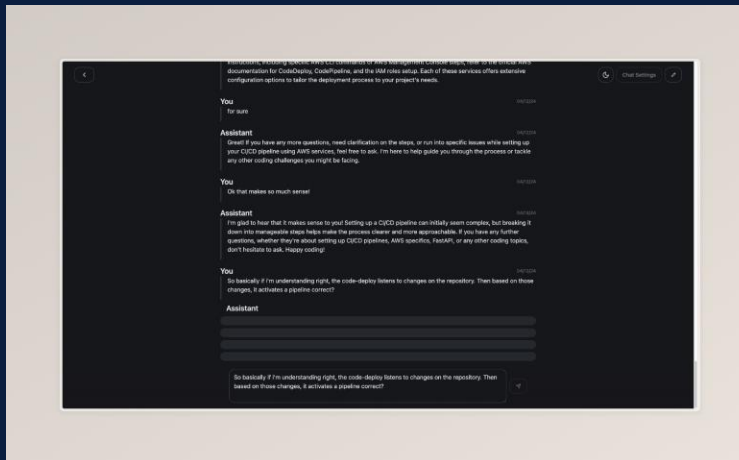
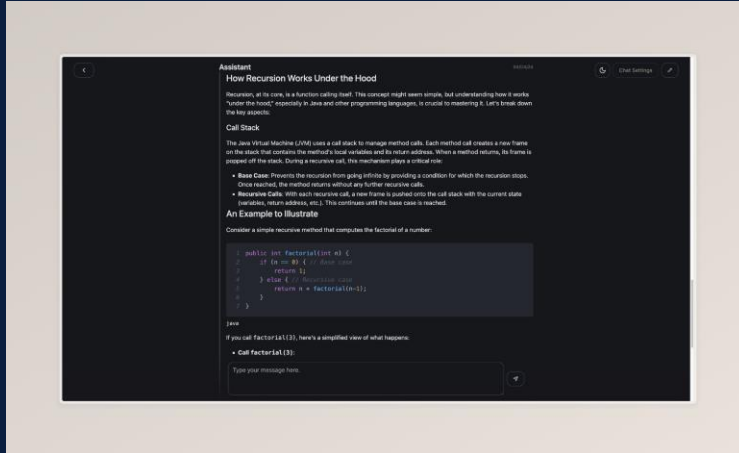
Using a PostgreSQL Database to store chat conversations, notes, and messages owned by the user.



Screenshots 1



Screenshots 2 (Dark Mode)





Screenshots (Backend)

```
INFO: Will watch for changes in these directories: ['/home/ubuntu/Backend']
INFO: Uvicorn running on http://0.0.0.0:10000 (Press CTRL+C to quit)
INFO: Started reloader process [13997] using WatchFiles
INFO: Started server process [14002]
INFO: Waiting for application startup.
INFO: Application startup complete.
INFO: 127.0.0.1:42432 - "OPTIONS /api/response?conversation_id=a6757d6e-560f-4999-b622-590564676de8&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:42444 - "POST /api/response?conversation_id=a6757d6e-560f-4999-b622-590564676de8&model=gpt-3.5-turbo HTTP/1.0" 200 OK
INFO: 127.0.0.1:33796 - "OPTIONS /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:33798 - "POST /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:38310 - "POST /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:53304 - "POST /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:35090 - "POST /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:42238 - "POST /api/response?conversation_id=363ce8d4-74bb-476c-9131-87c48feeab15&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
INFO: 127.0.0.1:37684 - "OPTIONS /api/response?conversation_id=9c7bf20b-551c-4536-a642-0218b1878e59&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:37692 - "POST /api/response?conversation_id=9c7bf20b-551c-4536-a642-0218b1878e59&model=gpt-3.5-turbo HTTP/1.0" 200 OK
INFO: 127.0.0.1:43136 - "OPTIONS /api/response?conversation_id=6df40d23-509f-464b-bfc2-bdd81960124a&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
gpt-4-turbo-preview
INFO: 127.0.0.1:43138 - "POST /api/response?conversation_id=6df40d23-509f-464b-bfc2-bdd81960124a&model=gpt-4-turbo-preview HTTP/1.0" 200 OK
INFO: 127.0.0.1:55280 - "OPTIONS /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:55282 - "POST /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:35862 - "POST /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:43946 - "POST /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:46984 - "POST /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
gpt-3.5-turbo
INFO: 127.0.0.1:38838 - "POST /api/response?conversation_id=2ab22add-cdbb-405d-8ca8-6a82112b66c9&model=gpt-3.5-turbo HTTP/1.0" 200 OK
```





Student Contributions

- Vincent Carrancho – Frontend/Backend
 - System Architecture and backend/frontend setup and development
- Robin Obregon – Backend
 - Backend Logic and deployment
- Stephanie Torres – Frontend
 - Conversation page and message loading/error states.
- Nicholas Belschner – Frontend
 - Account management pages and form schemas.
- Garvee Valcourt – Backend
 - Backend Logic with Robin and development.





Thank you!