



Meet the Founders of Anote



Natan Vidra

Data Scientist at Deloitte Applied AI

Electrical & Computer Engineering | Computer Science @
Cornell University

Deloitte.



Thomas Clifford

Software Engineer at Google for 5+ years

Financial Engineering | Computer Science @
Washington University

Google



Presentation Overview

Overview of Anote's Capabilities:

- A end to end MLOps platform that enables you to obtain the best large language model for your data.
- A evaluation framework to compare zero shot LLMs like GPT, Claude, Llama3 and Mistral, with fine tuned LLMs that are trained on your domain specific training data (via supervised, unsupervised and RLHF fine tuning).
- A data annotation interface to convert raw unstructured data into an LLM ready format, and incorporate subject matter expertise into your training process to improve model accuracies.
- An integration layer to route the best LLM into your own, on premise, private chatbot, via our fine tuning SDK.

Product Demos of Two Use Cases:

- Generative AI Red Teaming for Chatbots
- Automated CUI Tagging for Text Classification

Notable Customers

S&P Global

Marubeni

Weill Cornell

Harvard Medical School

US Navy

M2 Compliance



CHAT HISTORY



Document content info  

AI Engineering Studen  

Text Data Solutions Pl  

Chat 407  

Chat 374  

Chat 373  

Chat 372  

AI Startup Panacea So  

Chat 153  

Chat 152  

Healthcare Summariza  

Chat 19  

Answer Assistance Ser  



Text Data Solutions Platform



Hello, I am your financial assistant, upload a file to get started.

what is anote's solution?

Anote's solution consists of three fully functional products and three main steps:

- 1) Label Text Data - This entails classifying text, extracting entities, and answering questions on documents.
- 2) Fine Tune Model - Allows you to run your Fine Tuned LLM locally with their model inference API.
- 3) Private Chatbot - Enables you to chat with your documents while keeping your data private and secure. The core technology is a fine-tuning library that leverages state-of-the-art few-shot learning to make high-quality predictions with a few labeled samples. The fine-tuning can be done in three ways: 1) Unsupervised Fine Tuning -



Ask your document a question



TASK TYPES

File Uploader 

MODEL SELECTION

OpenAI 

Your own fine-tuned model key:

Model key

UPLOADED FILES

 PitchDeck_Anote.pdf 

 PrivateChatbotDeck.pdf

 Anote - Solution.pdf 

SOURCES

PitchDeck_Anote.pdf 

The Problems

Model Hallucinations:

As the volume of intelligence reports and classified documents increases, and questions become highly domain-specific, the model often hallucinates, leading to inaccurate or misleading responses.

Ex. "The military operation involved 23,700 personnel" vs. "The military operation involved 2,370 personnel."

Lack of Domain Specificity:

When asked questions about specific national security domains, such as intelligence gathering or defense strategies, the model's responses are often too vague or imprecise.

Ex. "'The strategy was effective in safeguarding national interests" vs. a detailed explanation of the defense measures used and their outcomes."

Adversarial Vulnerabilities:

AI systems used in national security are vulnerable to adversarial attacks such as data poisoning, where malicious data corrupts the training process, or evasion attacks, where inputs are altered to deceive the AI into making incorrect predictions.

Ex. "The AI system correctly classified this intelligence report" vs. "The AI system misclassified this report due to adversarial changes in the data."

Use Case 1 - Generative AI Red Teaming

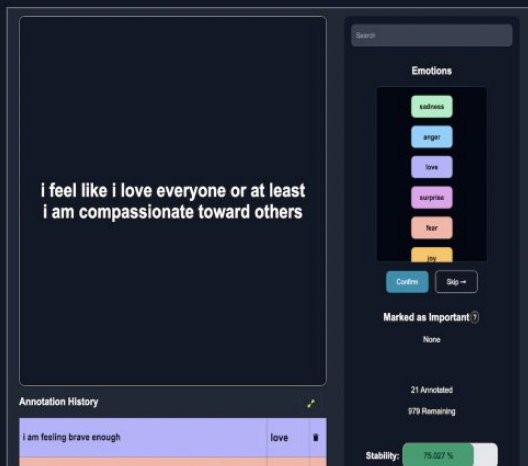
Anote co-led the U.S.-wide red-teaming challenge, NIST's ARIA pilot evaluation of LLM risks:

The first phase of this challenge tested and evaluated the robustness, security, and ethical implications of cutting-edge AI systems through adversarial testing. Participants identified exploits in three proxy scenarios, embodying issues of data exfiltration, bias, and hallucination.

The second phase of this challenge is an in-person exercise to be held alongside CAMLIS, that will include a red team evaluation using generative AI models. The in-person exercise will use the “Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1).”

<https://ai-challenges.nist.gov/aria>

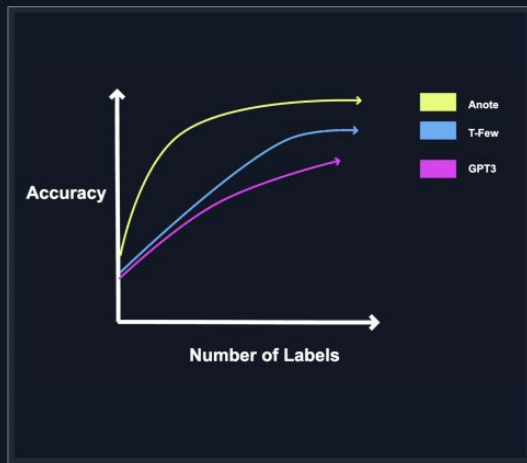
Our Solution



STEP 1

Label Text Data

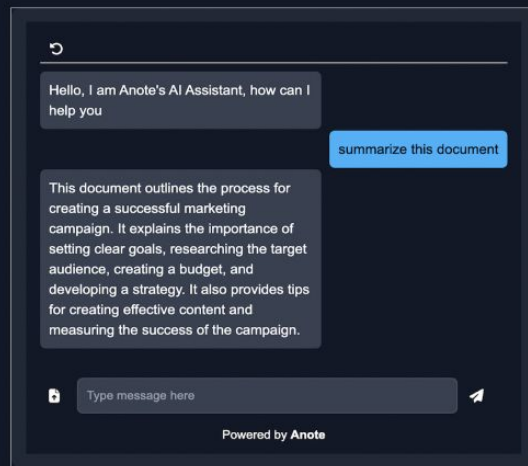
Classify Text, Extract Entities, and Answer Questions on Documents with LLMs



STEP 2

Fine Tune Model

Run your fine-tuned LLMs locally with our Model Inference API



STEP 3

Build Your Own Private Chatbot

Chat with your documents with LLMs while keeping your data private/secure

Data Labeler

State of the art few shot learning to make high quality predictions with a few labeled samples.



Supervised Fine-tuning

Fine-tune your model on your labeled data



Unsupervised Fine-tuning

Fine-tune your model from your raw unstructured documents

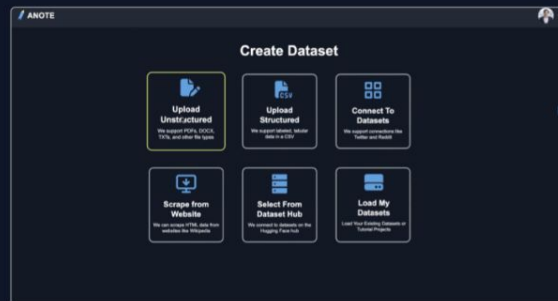


RLHF / RLAIF

Actively improve your models from human / AI feedback

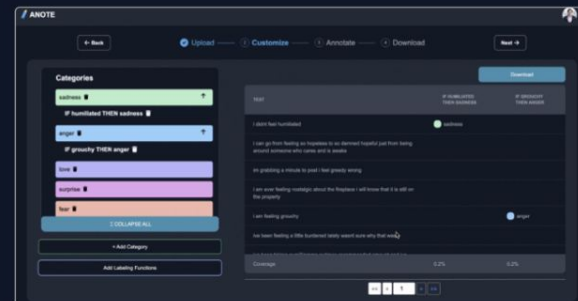
Upload

Create a new text based dataset.



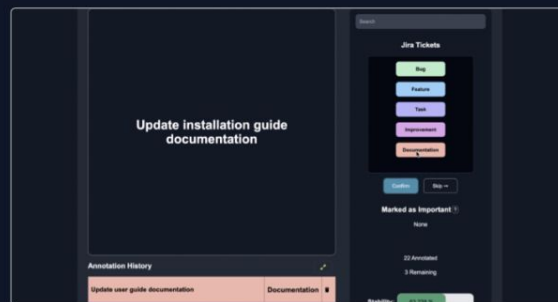
Customize

Add the categories, entities or questions you care about



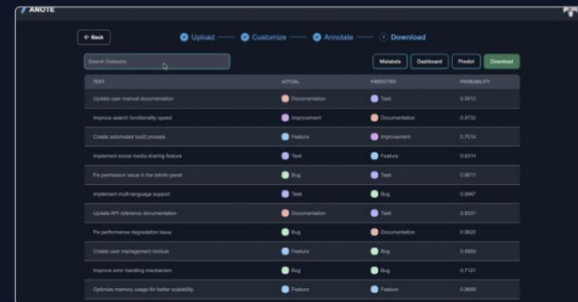
Annotate

As you annotate a few edge cases, the model actively learns to predict the rest.



Download

Download the resulting labels as a CSV. Export fine tuned model as an API endpoint



Private Chatbot

Your accurate private enterprise AI assistant



Upload

Upload your financial documents (10-Ks, 10-Qs, Earnings Calls Transcripts)



Chat

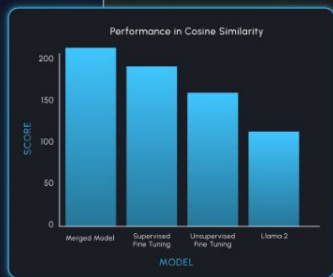
Ask Questions on your documents with models like GPT, Claude, Llama2 or Mistral



Evaluate

Get Citations for answers, and ensure the answer is correct to mitigate hallucinations

Evaluate



Input Fine-tune Model

Chat History

Jake's Budget Spending

Jake's Schedule

🔄

Jake's Budget Spending

📄

Hello I am Panacea, your personal AI assistant

What's my name and where am I from?

Your name is Jake and you grew up in Pennsylvania

How much did I spend in Feb 2023 versus Jan 2024?

You spent \$437 more in February 2023 than you did in January 2024.

Documents: 2023 Budget.pdf 2024 Budget.pdf
2023 Budget.pdf on page 2 line 37 expenses were \$3456 and 2024 Budget.pdf on page 4 line 20 where expenses were \$3019

📄 Ask your document a question

Model Selection

Jake's Fine-tuned Model

Uploaded Files

2023 Budget.pdf

2024 Budget.pdf

Upload Data

Sources

2023 Budget.pdf

Groceries \$250.29 Gym Membership \$80.98 Personal Furniture \$76.19 Rent \$1600.00 Supplies \$124.21 Transportation \$5.22 Utilities \$196.98 Total Expenses \$3456

2024 Budget.pdf

```
from privatechatbot import PrivateChatbot

api_key = "INSERT_API_KEY_HERE"
privatechatbot = PrivateChatbot(api_key)

chat_id = privatechatbot.upload(task_type="edgar", model_type="gpt", ticker="aapl")

response = privatechatbot.chat(chat_id, "What does this company do?")
print(response['answer'])
```

Software Development Kit

Evaluation Results

 Question Answering

 Structured

+ Add Dashboard View

Download Report

Models:

Fine-tuned Model

Claude

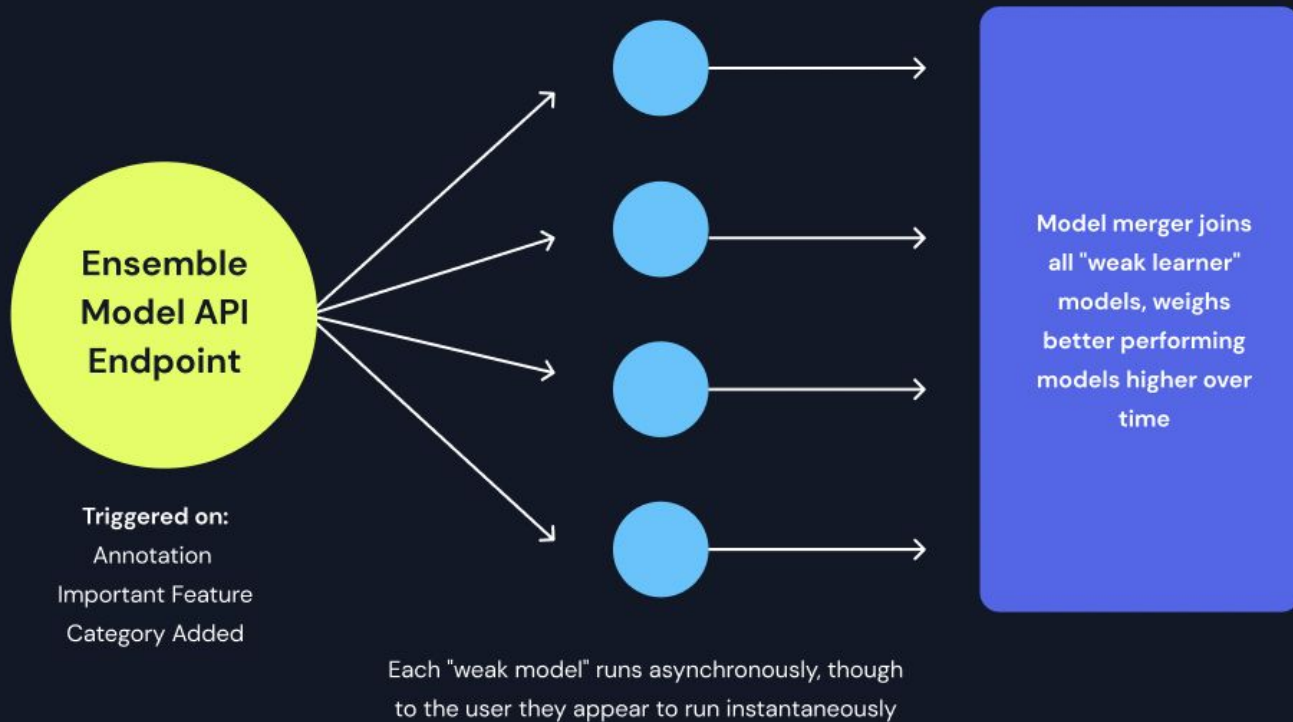
GPT 4

Evaluation Metric Scores ⓘ

Score	Fine-tuned Model	 Claude	 Open AI
Cosine Similarity Score	0.778	0.821	0.621
Rouge-L	0.824	0.901	0.780
LLM Evaluation Score	0.824	0.901	0.780

+ ADD A EVALUATION SCORE

the best model for your data



Differentiators



ACCURATE PREDICTIONS

Fine tuning and enhanced RAG for more accurate and tailored predictions. Our AI models actively learn and rapidly improve from SMEs.



ACCURATE CITATIONS

Accurate sources (page number, chunk of text, important features) to explain the models predictions and mitigate hallucinations.



COMPREHENSIVE CAPABILITIES

Supervised, Unsupervised and RLHF / RLAI fine tuning for classifying text, extracting entities, answering questions, and chatting with documents



EASY TO USE

Accessible UI similar to ChatGPT, and simple SDK for developers where you can input a fine tuned model for improved results.



PRIVATE VERSION

On premise enterprise-grade solution using Llama2 and GPT4All to leverage LLMs on your unstructured documents while keeping your data local and secure.



EVALUATION FRAMEWORK

Robust evaluation framework with metrics like Ragas Rouge-L, Cosine Similarity and Answer Relevance to show fine tuned model performance improvements

Why It Matters

Before

After Anote

Data Labeling

Manually Labeling their data themselves in a spreadsheet

State of the art few shot learning to make high quality predictions with a few labeled samples.

Tedious, Time Consuming, Costly

Less time, less expensive, higher accuracy

Manual Iterative Relabeling

Rapid flexibility for changing business requirements

Document Processing

Given a raw unstructured documents, such as a 10-k or earnings call transcript, you can't get answers to the questions right if trying to extract info, where accuracy really matters.

After a few interventions, we go from 10 questions right, to 15 questions right, to 20 questions right, to enable insights that were otherwise impossible to obtain.

Not accurate and largely manual extraction

Higher accuracy for raw unstructured documents

Sub-optimal analytics for critical business decisions

New insights that otherwise were not obtainable

Use Case 2 - The CUI Problem

The DOD generates millions of documents (PDFs, PPTX, DOCX, CSV, XLSX, EMAILS) per month that need to be tagged for Controlled Unclassified Information (CUI) content.

These documents contain sensitive information, so it is crucial to tag these documents with correct CUI categories to prevent highly classified documents from falling into the wrong people's hands.

<https://www.dodcui.mil/CUI-Registry-New/>

<https://www.challenge.gov/?challenge=cui>

The Problem with the Current Approach



Time Consuming



Requires Expertise



Not Accurate



Costly



Tedious



Not Scalable

Operational Resilience

If documents are under classified, people with not appropriate privilege levels have access to information that they are not allowed to access, which is a major security risk.

If documents are over classified, DOD personnel that should be able to collaborate and work on solving critical national security problems are not able to have access, which causes massive inefficiencies in productivity, communications, collaboration and getting things done.

Not solving the CUI problem is a national security risk, and will directly affect our ability to win future wars in the digital age, where information sharing is mission critical.

Time and Cost Savings

Team	DOD Current Process	Anote's Proposal
Number of People	80	5
Cost Per Person	\$70,000/year	\$70,000 / year
Cost of Product POC	\$0	\$X
Total Time to Label	6 months	10 days
Total Costs Per Month	\$2,800,000	\$175,000 + \$X

Requirements for a CUI Solution



Easy to Use



Private



Accurate



Citations



Evaluate



Recomposition

Our Solution

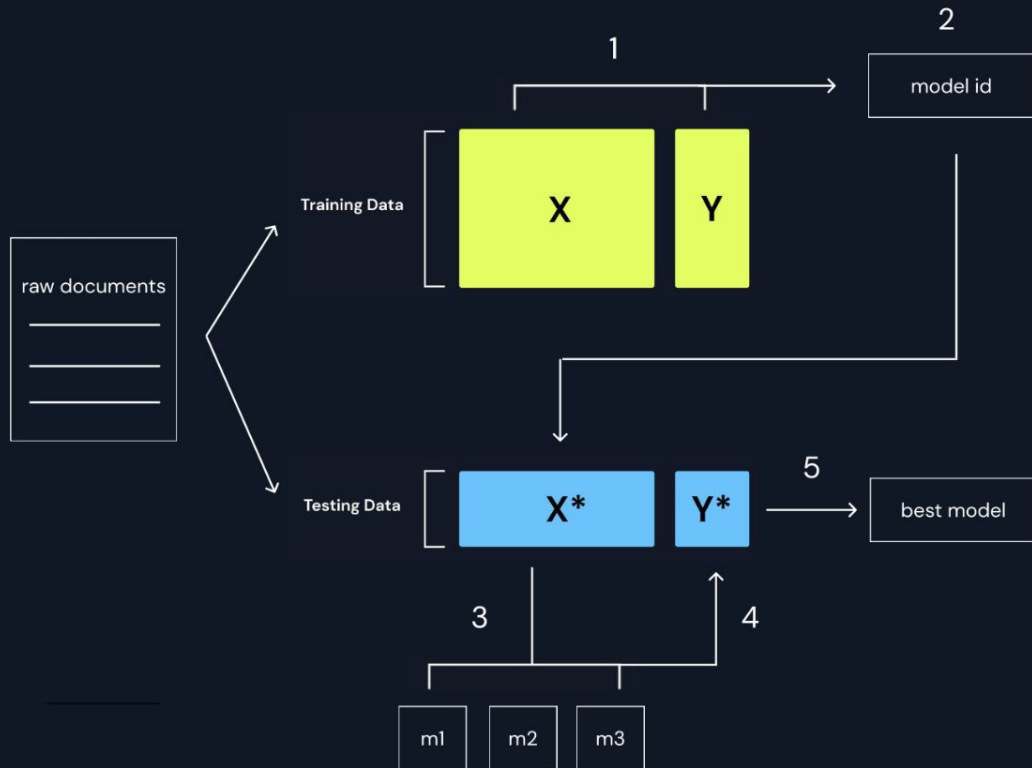
1. Label Data

2. Train Model

3. Make Predictions

4. Evaluate Results

5. Choose Best Model



Product Demo Overview

1

HOW TO ACCURATELY TAG CUI CONTENT

2

HOW TO IDENTIFY MISLABELED CUI TAGS

3

HOW TO SCALE A ROBUST CUI SOLUTION

Evaluation Results

Text Classification

Structured

+ Add a Dashboard View

DOWNLOAD REPORT

Models:

Ratex's Fine-tuned Model

Setfit

BERT

+

Classification Report Metrics

Model	MPC Accuracy	F1	Precision	Recall	Support
PROCURE	0.978	0.778	0.821	0.821	10
PRVCY	0.924	0.824	0.901	0.780	10
LEI	0.846	0.946	0.702	0.924	10
HLTH	0.945	0.776	0.785	0.924	10
Average/Total	0.987	0.876	0.965	0.824	10

Confusion Matrix

PROCURE	34	3	2	3
LEI	2	39	2	1
PRVCY	5	3	27	3
HLTH	2	1	2	42
	PROCURE	LEI	PRVCY	HLTH

Identifying Mislabels

Document Name	Human Label	Model Prediction
Doc1.txt	PRVCY	HLTH
Doc2.pdf	HLTH	LEI
Doc5.pptx	HLTH	PROCURE
Doc23.csv	PROCURE	LEI
Doc47.xlsx	LEI	HLTH

Statement of Work

Phase 1 - MVP

1. Predict categories on test documents from within the predict table view
2. Training the models on training documents to improve the performance via supervised fine tuning
3. Evaluate the models performance to see how different models perform on the testing documents
4. Baseline recomposition for standard PDFs, DOCx and PDF files

Phase 2 - Production Ready / Integration

1. Data labeler - incorporate subject matter expertise, and implement the RLHF training to make the models more accurate.
2. Core user permissions - projects, datasets, models, admin and annotator roles
3. Expand to more categories / questions and more documents
4. Integrate the fine tuned model with the Private Chatbot for on-premise deployment

Thank You!

Anote is a startup in New York City, helping make artificial intelligence more accessible. We believe there is a massive gap between the tremendous power of AI models, and the everyday tasks that people care about.

Contact Info



<https://anote.ai/>



nvidra@anote.ai



<https://www.linkedin.com/company/anote-ai/>



<https://www.youtube.com/@anote-ai>



[BENCHMARKING Q&A MODELS TALK](#)

Feb 2024 | Katherine Jijo



[BENCHMARKING TEXT CLASSIFICATION TALK](#)

Jan 2024 | Katherine Jijo



[FEW SHOT LEARNING TED TALK](#)

April 2023 | Natan Vidra



[IMPROVING RETRIEVAL FOR Q&A TALK](#)

Feb 2024 | Spurthi Setty



[FINE TUNING OF LLMS TALK](#)

Jan 2024 | Spurthi Setty